

Estimatori în procesarea semnalelor

Conf. dr. habil. Eduard Rothenstein

1 Considerente asupra mărginii inferioare Rao-Cramer (CRLB)

După cum am studiat la Metoda verosimilității maxime, determinarea unei margini inferioare pentru dispersia unui estimator nedeplasat se dovedește a fi foarte importantă în practică. **În cel mai bun caz**, se poate stabili dacă un estimator nedeplasat este de dispersie minimă (*estimator MVU - Minimum variance unbiased estimator*). **În cel mai rău caz**, furnizează un reper în raport cu care putem compara performanța oricărui alt estimator nedeplasat pentru parametrul estimat. Deși există metode alternative de determinare a unei margini inferioare pentru dispersie (vezi McAulay, Hofstetter [1971], Kendall, Stuart [1979], Seidman [1970], Ziv, Zakai [1969]), marginea inferioară Rao-Cramer (CRLB) este cea mai ușor de obținut. De asemenea, teoria ce conduce la stabilirea sa permite determinarea imediată a existenței unui estimator ce o atinge. Chiar dacă nu există un astfel de estimator, se pot determina estimatori ce ating acea margine într-un sens aproximant.

1.1 Acuratețea estimatorului

Cum toate informațiile privitoare la statistica analizată sunt încorporate în eșantionul observat și în densitatea de repartiție a caracteristicii, este evident că acuratețea estimării va depinde în mod direct de repartiția urmată. Este, prin urmare, clar că nu putem vorbi de acuratețea estimării unui parametru dacă densitatea nu depinde de acesta sau depinde în mod slab de el.

Exemplul 1.1 *Dependența densității de repartiție de parametrul necunoscut. Presupunem că, la o simplă observație, avem structura variabilei de selecție dată de*

$$X_0 = A + W_0, \quad \text{unde } W_0 \sim \mathcal{N}(0, \sigma^2)$$

și dorim să estimăm parametrul A . Desigur, este așteptată o estimare cu atât mai bună cu cât dispersia σ^2 este mai mică. Într-adevăr, un estimator nedeplasat bun este $\hat{A} = X_0$, varianța fiind σ^2 , deci acuratețea estimatorului crește odată cu descreșterea cantității σ^2 . Un mod alternativ, exemplificativ, de a observa aceasta este considerarea a două densități de repartiție, corespunzătoare la două dispersii $\sigma_1^2 < \sigma_2^2$:

$$f_i(x_0, A) = \frac{1}{\sqrt{2\pi\sigma_i^2}} \exp\left(-\frac{1}{2\sigma_i^2}(x_0 - A)^2\right), \quad i \in \{1, 2\} \quad (1)$$

Dacă reprezentăm cele două grafice pentru $x_0 = 3$ și $\sigma_1 = 1/3$, atunci se observă pe reprezentare că valori $A > 4$ sunt puternic improbabile. Pentru a vedea aceasta,

$$\mathbb{P}(3 - \delta/2 \leq X_0 \leq 3 + \delta/2) = \int_{3-\delta/2}^{3+\delta/2} f_i(t, A) dt,$$

care, pentru δ suficient de mic, este $f_i(x_0 = 3, A) \delta$. Dar

$$f_1(x_0 = 3, A = 4) \delta = 0.10\delta, \quad \text{în timp ce } f_1(x_0 = 3, A = 3) \delta = 1.2\delta.$$

Valorile parametrului $A > 4$ pot fi deci eliminate din analiză.

Aplicând teoria verosimilității maxime, considerăm logaritmul funcției de verosimilitate:

$$\ln f(x_0, A) = -\ln \sqrt{2\pi\sigma^2} - \frac{1}{2\sigma^2}(x_0 - A)^2, \quad \text{pentru care}$$

$$\frac{\partial \ln f(x_0, A)}{\partial A} = \frac{1}{\sigma^2}(x_0 - A) \quad \text{și} \quad -\frac{\partial^2 \ln f(x_0, A)}{\partial A^2} = \frac{1}{\sigma^2} \quad (2)$$

Curbura graficului crește pe măsură ce σ^2 descrește. Cum știam deja că estimatorul $\hat{A} = X_0$ are dispersia σ^2 , atunci, pentru acest exemplu,

$$D^2(\hat{A}) = \frac{1}{-\frac{\partial^2 \ln f(x_0, A)}{\partial A^2}} = \sigma^2.$$

Deși în acest exemplu, derivata secundă nu depinde de valoarea empirică observată x_0 , în general, va depinde. Din acest motiv, o măsură mai bună pentru curbura medie a graficului logaritmului funcției de verosimilitate este

$$-\mathbb{E}\left(\frac{\partial^2 \ln f(X_0, A)}{\partial A^2}\right). \quad (3)$$

Cu cât este mai mare cantitatea dată de (3), cu atât va fi mai mică dispersia estimatorului.

Revenind la aspectele teoretice prezentate la metoda verosimilității maxime, reamintim că, dacă densitatea de repartitie a caracteristicii studiate X este $f(x, \theta)$, atunci dispersia oricărui estimator nedeplasat $\hat{\theta}$ satisface relația

$$D^2(\hat{\theta}) \geq \frac{1}{-\mathbb{E}\left(\frac{\partial^2 \ln f(X, \theta)}{\partial \theta^2}\right)} = \frac{1}{I_1(\theta)} =: \frac{1}{I(\theta)} = CRLB(\hat{\theta}) \quad (4)$$

În plus, orice estimator nedeplasat care își atinge marginea (dispersiei) pentru toate valorile lui θ trebuie să satisfacă, ca și rezultat de caracterizare, relația:

$$\frac{\partial \ln f(x, \theta)}{\partial \theta} = I(\theta)(g(x) - \theta), \quad (5)$$

pentru funcțiile convenabil alese g și I . Acest estimator, care este un estimator MVU este de fapt

$$\hat{\theta} = g(X),$$

iar dispersia minimă atinsă este $1/I(\theta)$.

Exemplul 1.2 CRLB pentru Exemplul 1.1. Din formulele (2) și (4) obținem că

$$D^2(\hat{A}) \geq \sigma^2,$$

pentru toate valorile parametrului A . Prin urmare, nu poate exista niciun estimator nedeplasat a cărui dispersie să fie mai mică decât σ^2 , nici măcar pentru o singură valoare a parametrului A .

$$\text{Dar } \hat{A} = X_0 \text{ și atunci } D^2(\hat{A}) = \sigma^2.$$

Cum X_0 este nedeplasat și atinge CRLB, va fi un estimator MVU. Din (2) și (5) facem identificările:

$$\theta = A, \quad I(\theta) = \frac{1}{\sigma^2} \quad \text{și} \quad g(X_0) = X_0.$$

Prin urmare, $\hat{A} = g(X_0) = X_0$, deci este estimator MVU. De asemenea,

$$D^2(\hat{A}) = \sigma^2 = \frac{1}{I(\theta)}$$

și, în conformitate cu (4), avem

$$I(\theta) = -\mathbb{E}\left(\frac{\partial^2 \ln f(X, \theta)}{\partial \theta^2}\right).$$

Exemplul 1.3 Mărima (amplitudinea) curentului continuu în prezența zgomotului alb Gaussian (WGN, White Gaussian Noise). Generalizăm Exemplul 1.1 considerând cazul observațiilor multiple date de variabilele de selecție:

$$X_i = A + W_i, \quad i \in \{0, 1, \dots, n-1\}, \quad \text{unde } W_i \text{ este WGN, de dispersie } \sigma^2, \forall i$$

Pentru determinarea CRLB pentru parametrul A , considerăm funcția de verosimilitate

$$\begin{aligned} \mathcal{L}(V, A) = \mathcal{L}(X_0, \dots, X_{n-1}, A) &= \prod_{i=0}^{n-1} \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{1}{2\sigma^2}(X_i - A)^2\right) \\ &= \frac{1}{(2\pi\sigma^2)^{n/2}} \exp\left(-\frac{1}{2\sigma^2} \sum_{i=0}^{n-1} (X_i - A)^2\right), \end{aligned}$$

pentru care obținem,

$$\frac{\partial \ln \mathcal{L}(V, A)}{\partial A} = \frac{\partial}{\partial A} \left(-\ln \left((2\pi\sigma^2)^{n/2} \right) - \frac{1}{2\sigma^2} \sum_{i=0}^{n-1} (X_i - A)^2 \right) = \frac{1}{\sigma^2} \sum_{i=0}^{n-1} (X_i - A) = \frac{n}{\sigma^2} (\bar{X} - A), \quad (6)$$

iar derivata secundă este constantă

$$\frac{\partial^2 \ln \mathcal{L}(V, A)}{\partial \theta^2} = -\frac{n}{\sigma^2}.$$

De asemenea, din (6) rezultă că, anulând derivata de ordinul întâi, găsim estimatorul $\hat{A} = \bar{X}$.

Formula (4) conduce la următoarea cantitate pentru CRLB:

$$D^2(\hat{A}) = D^2(\bar{X}) \geq \frac{\sigma^2}{n} = CRLB(\hat{A}), \quad (7)$$

iar media de selecție este un estimator ce atinge acest prag minim, și deci, prin urmare, este un estimator MVU. El este, deci, și un estimator absolut corect.

Reamintim, de asemenea, că atunci când CRLB este atinsă, are loc egalitatea

$$D^2(\hat{\theta}) = \frac{1}{I(\theta)}, \quad \text{unde} \quad I(\theta) = -\mathbb{E} \left(\frac{\partial^2 \ln \mathcal{L}(V, \theta)}{\partial \theta^2} \right).$$

Următorul exemplu arată că, condiția de margine inferioară CRLB pentru dispersie nu este întotdeauna satisfăcută.

Exemplul 1.4 Estimarea fazei unui curent de tip sinusoidal într-un WGN. Considerăm problema determinării unui estimator pentru parametrul fazei ϕ , al unui curent sinusoidal, aflată într-un WGN:

$$X_i = A \cos(2\pi f_0 i + \phi) + W_i, \quad i \in \{0, 1, \dots, n-1\}. \quad (8)$$

Amplitudinea A și frecvența f_0 se presupun a fi cunoscute. Densitatea de repartiție a vectorului datelor este:

$$\mathcal{L}(V, \phi) = \frac{1}{(2\pi\sigma^2)^{n/2}} \exp \left(-\frac{1}{2\sigma^2} \sum_{i=0}^{n-1} (X_i - A \cos(2\pi f_0 i + \phi))^2 \right).$$

Derivatele logaritmului funcției de verosimilitate sunt:

$$\begin{cases} \frac{\partial \ln \mathcal{L}(V, \phi)}{\partial \phi} &= -\frac{1}{\sigma^2} \sum_{i=0}^{n-1} (X_i - A \cos(2\pi f_0 i + \phi)) A \sin(2\pi f_0 i + \phi) \\ &= -\frac{A}{\sigma^2} \sum_{i=0}^{n-1} \left(X_i \sin(2\pi f_0 i + \phi) - \frac{A}{2} \sin(4\pi f_0 i + 2\phi) \right) \\ \frac{\partial^2 \ln \mathcal{L}(V, \phi)}{\partial \phi^2} &= -\frac{A}{\sigma^2} \sum_{i=0}^{n-1} (X_i \cos(2\pi f_0 i + \phi) - A \cos(4\pi f_0 i + 2\phi)). \end{cases}$$

Având în vedere faptul că

$$\mathbb{E}(X_i) = \mathbb{E}(A \cos(2\pi f_0 i + \phi) + W_i) = A \cos(2\pi f_0 i + \phi) + \mathbb{E}(W_i) = A \cos(2\pi f_0 i + \phi),$$

și relația $\cos(2x) = 2\cos^2(x) - 1$, obținem astfel următoarea formulă pentru informația Fisher corespunzătoare parametrului ϕ :

$$\begin{aligned} -\mathbb{E} \left(\frac{\partial^2 \ln \mathcal{L}(V, \phi)}{\partial \phi^2} \right) &= \frac{A}{\sigma^2} \sum_{i=0}^{n-1} (A \cos^2(2\pi f_0 i + \phi) - A \cos(4\pi f_0 i + 2\phi)) \\ &= \frac{A^2}{\sigma^2} \sum_{i=0}^{n-1} \left(\frac{1}{2} + \frac{1}{2} \cos(4\pi f_0 i + 2\phi) - \cos(4\pi f_0 i + 2\phi) \right) \simeq \frac{nA^2}{2\sigma^2}, \end{aligned}$$

deoarece

$$\frac{1}{n} \sum_{i=0}^{n-1} \cos(4\pi f_0 i + 2\phi) \simeq 0 \quad \text{pentru} \quad f_0 \text{ neapropiat de } 0 \text{ sau } 1/2.$$

Prin urmare,

$$D^2(\hat{\phi}) > \frac{2\sigma^2}{nA^2}.$$

În acest exemplu, condiția pentru existența marginii inferioare atinse nu este satisfăcută. În consecință, **nu există un estimator pentru fază care să fie nedeplasat și care să atingă CRLB**.

Este posibil însă ca să existe un estimator de dispersie minimă, MVU. Nu putem, pentru moment, să stabilim dacă un estimator MVU pentru parametru există sau nu, iar dacă el există, cum să îl determinăm. Teoria Statisticilor suficiente ne va permite să găsim răspunsuri la aceste întrebări.

Așa cum am văzut, un estimator nedeplasat care atinge CRLB este un estimator eficient. Un estimator MVU nu este, neapărat, și eficient. Prin urmare, dacă marginea CRLB este atinsă, dispersia estimatorului este inversa informației Fisher. Intuitiv, cu cât avem la dispoziție mai multe informații, cu atât marginea inferioară a dispersiei estimatorului este mai mică. Pentru observații **ne-independente**, este de așteptat ca informația Fisher să fie mai mică decât $nI_1(\theta)$. Pentru observații **complet dependente**,

$$I_n(\theta) = I_1(\theta),$$

adică observații suplimentare nu vor aduce informații noi, ceea ce va face ca CRLB să nu descrească odată cu creșterea numărului de date noi observate.

Exemplul anterior poate fi extins la situația generală a determinării marginii inferioare a estimatorului $\hat{\theta}$ pentru semnale electrice de tipul următor, în care nu se dă forma explicită a dependenței semnalului de parametrul θ :

$$X_i = s_i(\theta) + W_i, \quad i \in \{0, 1, \dots, n-1\}.$$

Funcția de verosimilitate este, deoarece $X_i \sim \mathcal{N}(s_i(\theta), \sigma^2)$,

$$\mathcal{L}(V, \phi) = \frac{1}{(2\pi\sigma^2)^{n/2}} \exp\left(-\frac{1}{2\sigma^2} \sum_{i=0}^{n-1} (X_i - s_i(\theta))^2\right),$$

iar diferențialele logaritmului său sunt:

$$\begin{cases} \frac{\partial \ln \mathcal{L}(V, \phi)}{\partial \phi} = \frac{1}{\sigma^2} \sum_{i=0}^{n-1} (X_i - s_i(\theta)) \frac{\partial s_i(\theta)}{\partial \theta} \\ \frac{\partial^2 \ln \mathcal{L}(V, \phi)}{\partial \phi^2} = \frac{1}{\sigma^2} \sum_{i=0}^{n-1} \left((X_i - s_i(\theta)) \frac{\partial^2 s_i(\theta)}{\partial \theta^2} - \left(\frac{\partial s_i(\theta)}{\partial \theta} \right)^2 \right). \end{cases}$$

Aplicăm media și obținem, deoarece $\mathbb{E}(X_i - s_i(\theta)) = \mathbb{E}(W_i) = 0$, iar $\partial^2 s_i(\theta) / \partial \theta^2$ este deterministă,

$$\mathbb{E}\left(\frac{\partial^2 \ln \mathcal{L}(V, \phi)}{\partial \phi^2}\right) = -\frac{1}{\sigma^2} \sum_{i=0}^{n-1} \left(\frac{\partial s_i(\theta)}{\partial \theta}\right)^2.$$

În consecință,

$$D^2(\hat{\theta}) \geq \frac{\sigma^2}{\sum_{i=0}^{n-1} \left(\frac{\partial s_i(\theta)}{\partial \theta}\right)^2} = \frac{1}{I(\theta)}. \quad (9)$$

În particular, dacă semnalul curentului sinusoidal este reprezentată de (8), cu s_i depinzând de frecvența f_0 în modul următor:

$$s_i(f_0) := A \cos(2\pi f_0 i + \phi), \quad f_0 \in (0, 1/2),$$

amplitudinea A și faza ϕ fiind cunoscute, atunci estimatorul frecvenței curentului are marginea CRLB:

$$D^2(\hat{f}_0) \geq \frac{\sigma^2}{A^2 \sum_{i=0}^{n-1} (2\pi i \sin(2\pi f_0 i + \phi))^2} = CRLB(\hat{f}_0).$$

Menționăm că, dacă

$$f_0 \rightarrow 0, \quad \text{atunci} \quad CRLB(\hat{f}_0) \rightarrow +\infty.$$

Intuitiv, aceasta se datorează faptului că, pentru o frecvență f_0 suficient de mică (neglijabilă), o modificare a sa nu va altera semnalul într-o manieră semnificativă.

1.2 Transformarea parametrilor

De multe ori, în practică, se întâmplă ca parametrii pe care dorim să îi estimăm să depindă, sub forma unei funcții de alți parametri. De exemplu, putem fi interesați nu de estimarea nivelului curentului A , ci de puterea semnalului, A^2 . În general, dacă dorim să estimăm parametrul $\alpha = g(\theta)$, atunci *CRLB este*:

$$D^2(\hat{\alpha}) \geq \frac{\left(\frac{\partial g}{\partial \theta}\right)^2}{-\mathbb{E}\left(\frac{\partial^2 \ln \mathcal{L}(V, \theta)}{\partial \theta^2}\right)}. \quad (10)$$

Particularizând, în exemplul anterior, pentru $\alpha = g(A) = A^2$,

$$D^2(\hat{A^2}) \geq \frac{(2A)^2}{n/\sigma^2} = \frac{4A^2\sigma^2}{n} = CRLB(\hat{A^2}). \quad (11)$$

În Exemplul 1.3, media de selecție era un estimator eficient pentru A . Am putea presupune că $\bar{X}^2$ este, de asemenea, un estimator eficient și pentru A^2 . Pentru a evita această eroare de raționament arătăm, pentru început, că $\bar{X}^2$ nu este nici măcar un estimator nedeplasat. Deoarece

$$\bar{X} \sim \mathcal{N}(A, \sigma^2/n),$$

avem:

$$\mathbb{E}(\bar{X}^2) = \mathbb{E}^2(\bar{X}) + D^2(\bar{X}) = A^2 + \frac{\sigma^2}{n} \neq A^2. \quad (12)$$

Putem deci afirma imediat că **eficiența unui estimator este distrusă de o transformare neliniară a sa. Eficiența se păstrează însă în cazul transformărilor liniare.** Pentru a verifica această ultimă afirmație, presupunem că există un estimator eficient $\hat{\theta}$, pentru parametrul θ . Ne propunem acum să estimăm noul parametru

$$g(\theta) = a\theta + b, \quad \text{pentru } a, b \in \mathbb{R}.$$

Alegem drept estimator statistica

$$\widehat{g(\theta)} := g(\hat{\theta}) = a\hat{\theta} + b.$$

Avem astfel:

$$\begin{cases} \mathbb{E}(\widehat{g(\theta)}) &= \mathbb{E}(a\hat{\theta} + b) = a\mathbb{E}(\hat{\theta}) + b = a\theta + b = g(\theta), \quad \text{deci } \widehat{g(\theta)} \text{ este nedeplasat;} \\ D^2(\widehat{g(\theta)}) &\stackrel{(10)}{\geq} \frac{\left(\frac{\partial g}{\partial \theta}\right)^2}{I(\theta)} = \left(\frac{\partial g}{\partial \theta}\right)^2 D^2(\hat{\theta}) = a^2 D^2(\hat{\theta}) = D^2(a\hat{\theta} + b) = D^2(\widehat{g(\theta)}), \end{cases}$$

ceea ce implică faptul că are loc egalitate în ultima relație, adică *CRLB* este atinsă.

Deși eficiența se conservă în cazul transformărilor liniare ale parametrilor, ea este păstrată "aproximativ" și în cazul transformărilor neliniare dacă eșantionul considerat este suficient de mare. Aceasta are semnificație din punct de vedere practic deoarece se dorește frecvent estimarea unor funcții generale dependente de parametrul de bază. Pentru a vizualiza afirmația anterioară revenim la exemplul estimării lui A^2 cu ajutorul statisticii $\bar{X}^2$. Cu toate că aceasta este deplasată, din (12) se poate observa că, asimptotic, este de fapt un estimator nedeplasat (atunci când $n \rightarrow +\infty$). El este deci doar corect, nu și absolut corect. În plus, cum

$$\bar{X} \sim \mathcal{N}(A, \sigma^2/n),$$

putem evalua dispersia

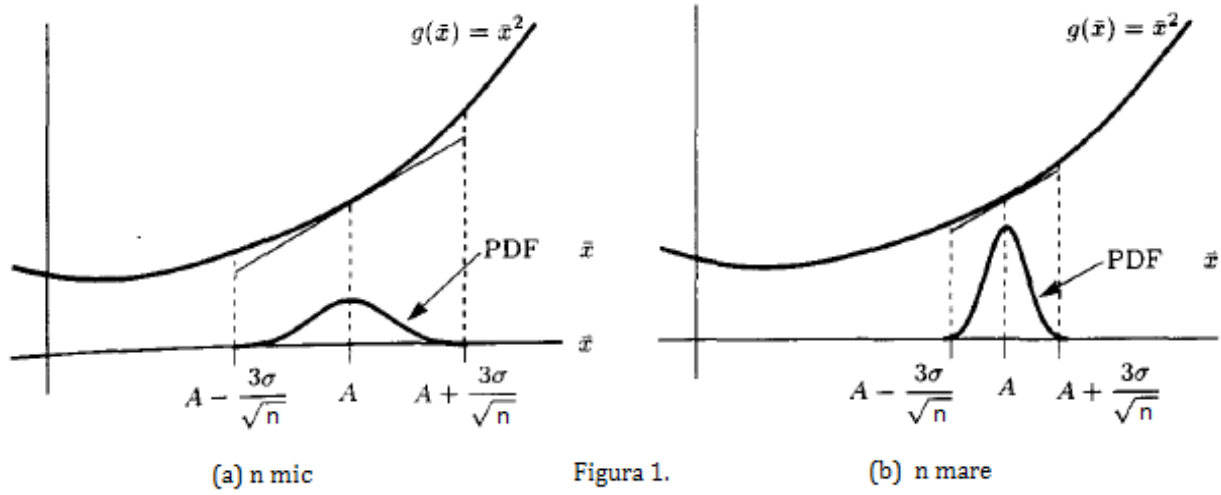
$$D^2(\bar{X}^2) = \mathbb{E}(\bar{X}^4) - \mathbb{E}^2(\bar{X}^2) = \frac{4A^2\sigma^2}{n} + \frac{2\sigma^4}{n^2}. \quad (13)$$

Pentru ultima egalitate, am folosit faptul că, dacă $\xi \sim \mathcal{N}(\mu, \sigma^2)$, atunci

$$\mathbb{E}(\xi^2) = \mu^2 + \sigma^2, \quad \mathbb{E}(\xi^4) = \mu^4 + 6\mu^2\sigma^2 + 3\sigma^4,$$

ceea ce conduce la

$$D^2(\bar{X}^2) = \mathbb{E}(\bar{X}^4) - \mathbb{E}^2(\bar{X}^2) = \mu^4 + 6\mu^2\sigma^2 + 3\sigma^4 - (\mu^4 + \sigma^4 + 2\mu^2\sigma^2) = 4\mu^2\sigma^2 + 2\sigma^4.$$



Se observă că al doilea termen din (13) converge la zero mai repede decât primul, ceea ce face ca, asimptotic, să regăsim CRLB din inegalitatea (11). Pentru un volum de selecție mare, media de selecție este concentrată în vecinătatea mediei teoretice A , iar pe acest interval mic transformarea neliniară este aproximativ liniară. De fapt, dacă aproximăm, prin liniarizare, funcția g în vecinătatea lui A , obținem:

$$g(\bar{x}) \simeq g(A) + g'(A)(\bar{x} - A), \quad \text{iar}$$

$$\mathbb{E}(g(\bar{X})) = g(A) = A^2, \quad \text{deoarece} \quad \mathbb{E}(\bar{X} - A) = \mathbb{E}(\bar{X}) - A = 0,$$

estimatorul fiind deci nedeplasat. De asemenea,

$$D^2(g(\bar{X})) = (g'(A))^2 D^2(\bar{X}) = \frac{(2A)^2 \sigma^2}{n} = \frac{4A^2 \sigma^2}{n},$$

adică estimatorul atinge, într-adevăr, **asimptotic**, CRLB.

1.3 Cazul parametrilor multipli (vectoriali)

Extindem analiza la cazul în care caracteristica studiată este dependentă de mai mulți parametri, vectorul acestora fiind notat cu $\theta = (\theta_1, \theta_2, \dots, \theta_p)^t$. Presupunem, de asemenea, că estimatorul $\hat{\theta}$ este nedeplasat. Marginea inferioară a dispersiei va fi un vector CRLB ce impune o limitare inferioară a fiecărei componente. Vom obține

$$\left| \begin{array}{l} D^2(\hat{\theta}_i) \geq (I^{-1}(\theta))_{ii}, \quad \text{unde} \quad I(\theta) \in \mathcal{M}_{p \times p}(\mathbb{R}) \text{ este matricea informației Fisher:} \\ (I(\theta))_{ij} = -\mathbb{E} \left(\frac{\partial^2 \ln \mathcal{L}(V, \theta)}{\partial \theta_i \partial \theta_j} \right), \quad \text{pentru } i, j \in \{1, 2, \dots, p\}. \end{array} \right. \quad (14)$$

Exemplul 1.5 Analiza curentului continuu bi-parametrizat în WGN. Revenim la Exemplul 1.3 și presupunem că, atât A , cât și σ^2 sunt necunoscute. Vectorul parametrilor este deci

$$\theta = (A, \sigma^2)^t, \quad p = 2.$$

Matricea informației Fisher asociate are forma:

$$I(\theta) = \begin{pmatrix} -\mathbb{E} \left(\frac{\partial^2 \ln \mathcal{L}(V, \theta)}{\partial A^2} \right) & -\mathbb{E} \left(\frac{\partial^2 \ln \mathcal{L}(V, \theta)}{\partial A \partial \sigma^2} \right) \\ -\mathbb{E} \left(\frac{\partial^2 \ln \mathcal{L}(V, \theta)}{\partial \sigma^2 \partial A} \right) & -\mathbb{E} \left(\frac{\partial^2 \ln \mathcal{L}(V, \theta)}{\partial (\sigma^2)^2} \right) \end{pmatrix}$$

Logaritmul funcției de verosimilitate și derivatele parțiale ale acestuia sunt:

$$\begin{aligned}\ln \mathcal{L}(V, \theta) &= -\frac{n}{2} \ln(2\pi) - \frac{n}{2} \ln \sigma^2 - \frac{1}{2\sigma^2} \sum_{i=0}^{n-1} (X_i - A)^2 \\ \frac{\partial \ln \mathcal{L}(V, \theta)}{\partial A} &= \frac{1}{\sigma^2} \sum_{i=0}^{n-1} (X_i - A), \quad \frac{\partial \ln \mathcal{L}(V, \theta)}{\partial \sigma^2} = -\frac{n}{2\sigma^2} + \frac{1}{2\sigma^4} \sum_{i=0}^{n-1} (X_i - A)^2, \quad \frac{\partial^2 \ln \mathcal{L}(V, \theta)}{\partial A^2} = -\frac{n}{\sigma^2}, \\ \frac{\partial^2 \ln \mathcal{L}(V, \theta)}{\partial A \partial \sigma^2} &= -\frac{1}{\sigma^4} \sum_{i=0}^{n-1} (X_i - A), \quad \frac{\partial^2 \ln \mathcal{L}(V, \theta)}{\partial (\sigma^2)^2} = \frac{n}{2\sigma^4} - \frac{1}{\sigma^6} \sum_{i=0}^{n-1} (X_i - A)^2.\end{aligned}$$

Matricea informației Fisher devine:

$$I(\theta) = \begin{pmatrix} \frac{n}{\sigma^2} & 0 \\ 0 & \frac{n}{2\sigma^4} \end{pmatrix}.$$

Pentru ultimul element din matricea anterioară am folosit evaluarea:

$$\begin{aligned}-\mathbb{E} \left(\frac{\partial^2 \ln \mathcal{L}(V, \theta)}{\partial (\sigma^2)^2} \right) &= -\mathbb{E} \left(\frac{n}{2\sigma^4} - \frac{1}{\sigma^6} \sum_{i=0}^{n-1} (X_i - A)^2 \right) = -\frac{n}{2\sigma^4} + \frac{1}{\sigma^6} \sum_{i=0}^{n-1} \mathbb{E} (X_i - A)^2 \\ &= -\frac{n}{2\sigma^4} + \frac{1}{\sigma^6} \sum_{i=0}^{n-1} \mathbb{E} (W_i^2) = -\frac{n}{2\sigma^4} + \frac{1}{\sigma^6} \sum_{i=0}^{n-1} (D^2(W_i) + \mathbb{E}^2(W_i)) \\ &= -\frac{n}{2\sigma^4} + \frac{1}{\sigma^6} \sum_{i=0}^{n-1} (\sigma^2 + 0) = -\frac{n}{2\sigma^4} + \frac{1}{\sigma^4} n = \frac{n}{2\sigma^4}.\end{aligned}$$

Din forma anterioară determinăm marginile

$$D^2(\hat{A}) \geq \frac{\sigma^2}{n} \quad \text{și} \quad D^2(\hat{\sigma}^2) \geq \frac{2\sigma^4}{n}.$$

Observația 1.1 Chiar dacă, pentru acest exemplu, matricea $I(\theta)$ este de tip diagonal, în cazul general ea nu are această proprietate. Remarcăm, de asemenea, că CRLB pentru estimatorul $\hat{A}$ este același ca și cazul în care σ^2 era cunoscut. Acest lucru se întâmplă doar datorită formei diagonale a matricei $I(\theta)$. În absența acestei proprietăți, nici identificarea anterioară nu mai este posibilă, după cum vedem în exemplul următor.

Exemplul 1.6 Aproximarea drepte de regresie (line fitting). Considerăm observațiile date de următoarele formule de dependență, ce vor conduce la estimarea drepte de regresie:

$$X_i = A + Bi + W_i, \quad \text{unde } W_i \text{ sunt WGN independente, } i \in \{0, 1, \dots, n-1\}.$$

Ne propunem să determinăm CRLB pentru parametrii A (intercept) și B (slope), vectorul parametrilor fiind $\theta = (A, B)^t$. Funcția de verosimilitate și derivatele parțiale ale logaritmului său sunt, deoarece $X_i \sim \mathcal{N}(A + Bi, \sigma^2)$:

$$\begin{aligned}\mathcal{L}(V, \theta) &= \frac{1}{(2\pi\sigma^2)^{n/2}} \exp \left(-\frac{1}{2\sigma^2} \sum_{i=0}^{n-1} (X_i - A - Bi)^2 \right) \\ \frac{\partial \ln \mathcal{L}(V, \theta)}{\partial A} &= \frac{1}{\sigma^2} \sum_{i=0}^{n-1} (X_i - A - Bi), \quad \frac{\partial \ln \mathcal{L}(V, \theta)}{\partial B} = \frac{1}{\sigma^2} \sum_{i=0}^{n-1} (X_i - A - Bi) i, \quad \frac{\partial^2 \ln \mathcal{L}(V, \theta)}{\partial A^2} = -\frac{n}{\sigma^2}, \\ \frac{\partial^2 \ln \mathcal{L}(V, \theta)}{\partial A \partial B} &= -\frac{1}{\sigma^2} \sum_{i=0}^{n-1} i, \quad \frac{\partial^2 \ln \mathcal{L}(V, \theta)}{\partial B^2} = -\frac{1}{\sigma^2} \sum_{i=0}^{n-1} i^2.\end{aligned}$$

Matricea informației Fisher are următoarea formă imediată, deoarece derivatele secunde nu depind de variabilele X_i :

$$I(\theta) = \frac{1}{\sigma^2} \begin{pmatrix} n & \sum_{i=0}^{n-1} i \\ \sum_{i=0}^{n-1} i & \sum_{i=0}^{n-1} i^2 \end{pmatrix} = \frac{1}{\sigma^2} \begin{pmatrix} n & \frac{n(n-1)}{2} \\ \frac{n(n-1)}{2} & \frac{n(n-1)(2n-1)}{6} \end{pmatrix}.$$

Inversa acestei matrici este:

$$I^{-1}(\theta) = \sigma^2 \begin{pmatrix} \frac{2(2n-1)}{n(n+1)} & -\frac{6}{n(n+1)} \\ -\frac{6}{n(n+1)} & \frac{12}{n(n^2-1)} \end{pmatrix}.$$

Formula (14) arată că CRLB sunt date de inegalitățile

$$D^2(\hat{A}) \geq \frac{2(2n-1)\sigma^2}{n(n+1)} = CRLB(\hat{A}) \quad \text{și} \quad D^2(\hat{B}) \geq \frac{12\sigma^2}{n(n^2-1)} = CRLB(\hat{B}).$$

Remarcăm faptul că CRLB pentru estimatorul $\hat{A}$ crește față de situația în care parametrul B era cunoscut, deoarece, pentru $n \geq 2$,

$$D^2(\hat{A}) \geq \frac{2(2n-1)}{n+1} \cdot \frac{\sigma^2}{n} \geq 1 \cdot \frac{\sigma^2}{n} = -\frac{1}{\mathbb{E}\left(\frac{\partial^2 \ln \mathcal{L}(V, A)}{\partial A^2}\right)},$$

ultimul termen identificându-se chiar cu CRLB obținută pentru $\hat{A}$, când B era cunoscut. Anticipăm astfel un rezultat general valabil, și anume faptul CRLB întotdeauna crește odată cu creșterea numărului de parametri estimați.

În altă ordine de idei,

$$\frac{CRLB(\hat{A})}{CRLB(\hat{B})} = \frac{\frac{2(2n-1)\sigma^2}{n(n+1)}}{\frac{12\sigma^2}{n(n^2-1)}} = \frac{(2n-1)(n-1)}{6} \geq 1, \quad \text{pentru } n \geq 3.$$

Prin urmare, parametrul B este mai ușor de estimat, marginea sa inferioară CRLB având o descreștere de ordinul lui $1/n^3$, în comparație cu dependența de ordinul $1/n$ a marginii CRLB a celuilalt parametru. **Această diferență de magnitudine indică faptul că observația X_i este mai sensibilă la modificări ale parametrului B față de modificări ale parametrului A .** Observăm, de asemenea, că

$$\Delta X_i \simeq \frac{\partial X_i}{\partial A} \Delta A = \Delta A \quad \text{și} \quad \Delta X_i \simeq \frac{\partial X_i}{\partial B} \Delta B = i \Delta B,$$

adică modificările lui B se resimt multiplicat în evoluția observațiilor (aleatoare) X_i .

Prezentăm acum varianta vectorială a Teoremei Rao-Cramer pentru marginea inferioară.

Teorema 1 Presupunem că densitatea de repartiție a vectorului aleator de selecție $\mathcal{L}(V, \theta)$ satisface următoarele condiții de regularitate:

$$\mathbb{E}\left(\frac{\partial \ln \mathcal{L}(V, \theta)}{\partial \theta}\right) = 0, \quad \text{pentru toți parametrii } \theta.$$

Atunci, matricea de covarianță $C_{\hat{\theta}}$ a oricărui estimator nedeplasat $\hat{\theta}$ satisface inegalitatea

$$C_{\hat{\theta}} - I^{-1}(\hat{\theta}) \geq \mathbf{0}, \quad (15)$$

inegalitatea fiind înțeleasă în sensul că matricea din membrul stâng este pozitiv semi-definită. Informația Fisher este matricea

$$(I(\theta))_{ij} = -\mathbb{E}\left(\frac{\partial^2 \ln \mathcal{L}(V, \theta)}{\partial \theta_i \partial \theta_j}\right).$$

În plus, un estimator nedeplasat atinge marginea inferioară în sensul că

$$C_{\hat{\theta}} = I^{-1}(\hat{\theta})$$

dacă și numai dacă

$$\frac{\partial \ln \mathcal{L}(V, \theta)}{\partial \theta} = I(\theta) (g(V) - \theta), \quad (16)$$

pentru o funcție $g: \mathbb{R}^n \rightarrow \mathbb{R}^p$ și o matrice $I \in \mathcal{M}_{p \times p}$.

Estimatorul, care este un estimator de dispersie minimă (MVU) este:

$$\hat{\theta} = g(V),$$

iar matricea sa de covarianță este $I^{-1}(\theta)$.

Demonstrație. Considerăm vectorul parametrilor $\alpha = g(\theta)$, densitatea de repartiție a caracteristicii fiind dependentă de vectorul θ al parametrilor. Fie estimatorii nedeplasați astfel încât

$$\mathbb{E}(\hat{\alpha}_i) = \alpha_i = (g(\theta))_i, \quad i \in \{1, 2, \dots, r\}.$$

Condițiile de regularitate conduc la

$$\int_{\mathbb{R}^n} (\hat{\alpha}_i - \alpha_i) \frac{\partial \ln \mathcal{L}(v, \theta)}{\partial \theta_i} \mathcal{L}(v, \theta) dv = \frac{\partial (g(\theta))_i}{\partial \theta_i}, \quad i \in \{1, 2, \dots, p\}. \quad (17)$$

Pentru $j \neq i$,

$$\begin{aligned} \int_{\mathbb{R}^n} (\hat{\alpha}_i - \alpha_i) \frac{\partial \ln \mathcal{L}(v, \theta)}{\partial \theta_j} \mathcal{L}(v, \theta) dv &= \int_{\mathbb{R}^n} (\hat{\alpha}_i - \alpha_i) \frac{\partial \mathcal{L}(v, \theta)}{\partial \theta_j} dv \\ &= \frac{\partial}{\partial \theta_j} \int_{\mathbb{R}^n} \hat{\alpha}_i \mathcal{L}(v, \theta) dv - \underbrace{\alpha_i \mathbb{E} \left(\frac{\partial \ln \mathcal{L}(v, \theta)}{\partial \theta_j} \right)}_{=0} = \frac{\partial \alpha_i}{\partial \theta_j} = \frac{\partial (g(\theta))_i}{\partial \theta_j} \end{aligned} \quad (18a)$$

Scriind (17) și (18a) sub formă matriceală obținem

$$\int_{\mathbb{R}^n} (\hat{\alpha} - \alpha) \frac{\partial \ln \mathcal{L}(v, \theta)}{\partial \theta} \mathcal{L}(v, \theta) dv = \frac{\partial g(\theta)}{\partial \theta}.$$

Considerăm acum vectorii, posibili arbitrari, $\mathbf{a} \in \mathcal{M}_{r \times 1}$ și $\mathbf{b} \in \mathcal{M}_{p \times 1}$. Multiplicând în pozițiile convenabile (la stânga / la dreapta) cu vectorii $\mathbf{a}^t$ și $\mathbf{b}$ și ținând cont de faptul că $\mathcal{L}(v, \theta)$ este un scalar, găsim:

$$\int_{\mathbb{R}^n} \underbrace{\mathbf{a}^t (\hat{\alpha} - \alpha)}_{\mathbf{a}^t (\hat{\alpha} - \alpha)} \cdot \underbrace{\frac{\partial \ln \mathcal{L}(v, \theta)}{\partial \theta}}_{\frac{\partial \ln \mathcal{L}(v, \theta)}{\partial \theta}} \cdot \underbrace{\mathcal{L}(v, \theta)}_{\mathcal{L}(v, \theta)} dv = \mathbf{a}^t \frac{\partial g(\theta)}{\partial \theta} \mathbf{b}. \quad (19)$$

Notăm

$$w(v) = \mathcal{L}(v, \theta), \quad g(v) = \mathbf{a}^t (\hat{\alpha} - \alpha), \quad h(v) = \frac{\partial \ln \mathcal{L}(v, \theta)}{\partial \theta} \mathbf{b}.$$

Ridicăm la pătrat (19) și, aplicând inegalitatea Cauchy Schwarz, obținem că:

$$\begin{aligned} \left(\mathbf{a}^t \frac{\partial g(\theta)}{\partial \theta} \mathbf{b} \right)^2 &\leq \int_{\mathbb{R}^n} \left(\sqrt{\mathcal{L}(v, \theta)} \mathbf{a}^t (\hat{\alpha} - \alpha) \right) \cdot \left(\sqrt{\mathcal{L}(v, \theta)} \frac{\partial \ln \mathcal{L}(v, \theta)}{\partial \theta} \mathbf{b} \right) dv \\ &\leq \int_{\mathbb{R}^n} \mathcal{L}(v, \theta) \mathbf{a}^t (\hat{\alpha} - \alpha) (\hat{\alpha} - \alpha)^t \mathbf{a} dv \cdot \int_{\mathbb{R}^n} \mathcal{L}(v, \theta) \mathbf{b}^t \frac{\partial \ln \mathcal{L}(v, \theta)}{\partial \theta} \frac{\partial \ln \mathcal{L}(v, \theta)}{\partial \theta} \mathbf{b} dv \\ &= \mathbf{a}^t C_{\hat{\alpha}} \mathbf{a} \cdot \mathbf{b}^t I(\theta) \mathbf{b}, \end{aligned} \quad (20)$$

deoarece, similar cazului scalar,

$$\mathbb{E} \left(\frac{\partial \ln \mathcal{L}(v, \theta)}{\partial \theta_i} \frac{\partial \ln \mathcal{L}(v, \theta)}{\partial \theta_j} \right) = -\mathbb{E} \left(\frac{\partial^2 \ln \mathcal{L}(v, \theta)}{\partial \theta_i \partial \theta_j} \right) = (I(\theta))_{ij}.$$

Cum $\mathbf{b}$ este arbitrar ales, alegem pentru el valoarea

$$\mathbf{b} := I^{-1}(\theta) \frac{\partial g(\theta)}{\partial \theta} \mathbf{a} \quad \text{deci} \quad \mathbf{b}^t = \mathbf{a}^t \frac{\partial g(\theta)}{\partial \theta} (I^{-1}(\theta))^t = \mathbf{a}^t \frac{\partial g(\theta)}{\partial \theta} I^{-1}(\theta).$$

Obținem din (20),

$$\begin{aligned} \left(\mathbf{a}^t \frac{\partial g(\theta)}{\partial \theta} I^{-1}(\theta) \frac{\partial g(\theta)}{\partial \theta} \mathbf{a} \right)^2 &\leq \mathbf{a}^t C_{\hat{\alpha}} \mathbf{a} \cdot \left(\mathbf{a}^t \frac{\partial g(\theta)}{\partial \theta} I^{-1}(\theta) I(\theta) I^{-1}(\theta) \frac{\partial g(\theta)}{\partial \theta} \mathbf{a} \right) \\ &= \mathbf{a}^t C_{\hat{\alpha}} \mathbf{a} \cdot \left(\mathbf{a}^t \frac{\partial g(\theta)}{\partial \theta} I^{-1}(\theta) \frac{\partial g(\theta)}{\partial \theta} \mathbf{a} \right). \end{aligned}$$

Deoarece $I(\theta)$ este pozitiv definită, aceeași proprietate o are și $I^{-1}(\theta)$, iar matricea

$$\frac{\partial g(\theta)}{\partial \theta} I^{-1}(\theta) \frac{\partial g(\theta)}{\partial \theta}^t$$

este, cel puțin, pozitiv semi-definită. Prin urmare, termenul din parantezele anterioare este ne-negativ și deducem că

$$\mathbf{a}^t \left(C_{\hat{\alpha}} - \frac{\partial g(\theta)}{\partial \theta} I^{-1}(\theta) \frac{\partial g(\theta)^t}{\partial \theta} \right) \mathbf{a} \geq 0.$$

Reamintim că și $\mathbf{a}$ este ales arbitrar, ceea ce conduce la

$$C_{\hat{\alpha}} - \frac{\partial g(\theta)}{\partial \theta} I^{-1}(\theta) \frac{\partial g(\theta)^t}{\partial \theta} \geq 0.$$

Dacă

$$\alpha = g(\theta) = \theta, \quad \text{atunci} \quad \frac{\partial g(\theta)}{\partial \theta} = I,$$

iar inegalitatea (15) este satisfăcută. Condiția pentru egalitate în inegalitatea Cauchy Schwarz este

$$g(v) = c \cdot h(v) \iff \mathbf{a}^t(\hat{\alpha} - \alpha) = c \frac{\partial \ln \mathcal{L}(v, \theta)^t}{\partial \theta} \mathbf{b},$$

unde c este o constantă independentă de v . Această condiție devine:

$$\mathbf{a}^t(\hat{\alpha} - \alpha) = c \frac{\partial \ln \mathcal{L}(v, \theta)^t}{\partial \theta} \mathbf{b} = c \frac{\partial \ln \mathcal{L}(v, \theta)^t}{\partial \theta} \underbrace{I^{-1}(\theta) \frac{\partial g(\theta)^t}{\partial \theta}}_{=\mathbf{b}} \mathbf{a}.$$

Cum vectorul $\mathbf{a}$ este arbitrar ales,

$$\frac{\partial g(\theta)}{\partial \theta} I^{-1}(\theta) \frac{\partial \ln \mathcal{L}(v, \theta)}{\partial \theta} = \frac{1}{c} (\hat{\alpha} - \alpha).$$

Dacă considerăm $\alpha = g(\theta) = \theta$, atunci

$$\frac{\partial \ln \mathcal{L}(v, \theta)}{\partial \theta} = \frac{1}{c} I(\theta) (\hat{\theta} - \theta).$$

Cum c poate depinde de vectorul parametrilor θ ,

$$\begin{aligned} \frac{\partial \ln \mathcal{L}(v, \theta)}{\partial \theta_i} &= \sum_{k=1}^p \frac{(I(\theta))_{ik}}{c(\theta)} (\hat{\theta}_k - \theta_k), \quad \text{de unde, diferențiind încă o dată,} \\ \frac{\partial^2 \ln \mathcal{L}(v, \theta)}{\partial \theta_i \partial \theta_j} &= \sum_{k=1}^p \left(\frac{(I(\theta))_{ik}}{c(\theta)} (-\delta_{kj}) + \frac{\partial \left(\frac{(I(\theta))_{ik}}{c(\theta)} \right)}{\partial \theta_j} (\hat{\theta}_k - \theta_k) \right). \end{aligned}$$

În final,

$$(I(\theta))_{ij} = -\mathbb{E} \left(\frac{\partial^2 \ln \mathcal{L}(v, \theta)}{\partial \theta_i \partial \theta_j} \right) = \frac{(I(\theta))_{ij}}{c(\theta)}, \quad \text{deoarece} \quad \mathbb{E}(\hat{\theta}_k) = \theta_k.$$

Evident,

$$c(\theta) = 1,$$

iar condiția pentru egalitate rezultă imediat. Demonstrația este, în acest moment, încheiată.

Faptul că relația (14) rezultă din (15) se justifică remarcând că pentru o matrice pozitiv semi-definită elementele diagonale sunt ne-negative. Prin urmare,

$$(C_{\hat{\theta}} - I^{-1}(\theta))_{ii} \geq 0$$

și, în consecință,

$$D^2(\hat{\theta}_i) \geq (C_{\hat{\theta}})_{ii} \geq (I^{-1}(\theta))_{ii}, \quad \text{pentru orice } i. \quad (21)$$

Când are loc egalitatea (adică $C_{\hat{\theta}} = I^{-1}(\theta)$), atunci și (21) are loc cu egalitate. Estimatorul

$$\hat{\theta} = g(V)$$

este eficient și, prin urmare, este un estimator MVU . Un caz privind egalitatea am obținut în Exemplul 1.6. Am determinat

$$\frac{\partial \ln \mathcal{L}(v, \theta)}{\partial \theta} = \begin{pmatrix} \frac{\partial \ln \mathcal{L}(v, \theta)}{\partial A} \\ \frac{\partial \ln \mathcal{L}(v, \theta)}{\partial B} \end{pmatrix} = \begin{pmatrix} \frac{1}{\sigma^2} \sum_{i=0}^{n-1} (X_i - A - Bi) \\ \frac{1}{\sigma^2} \sum_{i=0}^{n-1} (X_i - A - Bi) i \end{pmatrix}. \quad (22)$$

Putem scrie (22) sub caracterizarea echivalentă (se verifică prin calcul direct structura următoare)

$$\frac{\partial \ln \mathcal{L}(v, \theta)}{\partial \theta} = \begin{pmatrix} \frac{n}{\sigma^2} & \frac{n(n-1)}{2\sigma^2} \\ \frac{n(n-1)}{2\sigma^2} & \frac{n(n-1)(2n-1)}{6\sigma^2} \end{pmatrix} \begin{pmatrix} \hat{A} - A \\ \hat{B} - B \end{pmatrix}, \quad (23)$$

unde

$$\hat{A} = \frac{2(2n-1)}{n(n+1)} \sum_{i=0}^{n-1} X_i - \frac{6}{n(n+1)} \sum_{i=0}^{n-1} i X_i \quad \text{și} \quad \hat{B} = -\frac{6}{n(n+1)} \sum_{i=0}^{n-1} X_i + \frac{12}{n(n^2-1)} \sum_{i=0}^{n-1} i X_i.$$

Condițiile pentru egalitate sunt satisfăcute, iar estimatorul

$$\hat{\theta} = (\hat{A}, \hat{B})^t$$

este eficient și, prin urmare, este un estimator MVU . Matricea din (23) este inversa matricei de covarianță.

Dacă condiția de egalitate are loc, atunci, întotdeauna, estimatorul $\hat{\theta}$ este nedeplasat. Într-adevăr, condiția de regularitate, aplicată reprezentării (16) implică

$$\mathbb{E}(g(V)) = \mathbb{E}(\hat{\theta}) = \theta.$$

Pentru determinarea estimatorilor MVU pentru un parametru vectorial, Teorema CRLB este un instrument recomandat. În general, metoda este folosită în aplicații practice, pentru modele liniare, Exemplul 1.6 al estimării drepte de regresie (*line fitting*) fiind doar un caz particular.

1.4 Transformarea parametrilor vectoriali

Dorim să estimăm parametrul vectorial $\alpha = g(\theta)$, g fiind o funcție r -dimensională, $g: \mathbb{R}^p \rightarrow \mathbb{R}^r$. După cum am demonstrat în Teorema 1,

$$C_{\hat{\alpha}} - \frac{\partial g(\theta)}{\partial \theta} I^{-1}(\theta) \frac{\partial g(\theta)^t}{\partial \theta} \geq 0. \quad (24)$$

Prin $\partial g(\theta) / \partial \theta$ am înțeles, desigur, matricea Jacobiană $r \times p$ dimensională

$$\frac{\partial g(\theta)}{\partial \theta} = \begin{pmatrix} \frac{\partial g_1(\theta)}{\partial \theta_1} & \frac{\partial g_1(\theta)}{\partial \theta_2} & \dots & \frac{\partial g_1(\theta)}{\partial \theta_p} \\ \frac{\partial g_2(\theta)}{\partial \theta_1} & \frac{\partial g_2(\theta)}{\partial \theta_2} & \dots & \frac{\partial g_2(\theta)}{\partial \theta_p} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{\partial g_r(\theta)}{\partial \theta_1} & \frac{\partial g_r(\theta)}{\partial \theta_2} & \dots & \frac{\partial g_r(\theta)}{\partial \theta_p} \end{pmatrix}.$$

Considerăm acum un curent continuu într-un WGN, ce are parametrii A și σ^2 necunoscuți și ne propunem să estimăm parametrul

$$\alpha := \frac{A^2}{\sigma^2},$$

care poate fi considerat proporția semnal-zgomot (SNR , *Signal-to-Noise Ratio*). Vectorul parametrilor este

$$\theta = (\mathbf{A}, \sigma^2)^t, \quad \text{iar} \quad g(\theta) = \frac{\theta_1^2}{\theta_2} = \frac{A^2}{\sigma^2}.$$

Determinăm, prin calcul imediat:

$$\left\{ \begin{array}{l} I(\theta) = \begin{pmatrix} \frac{n}{\sigma^2} & 0 \\ 0 & \frac{n}{2\sigma^4} \end{pmatrix}, \text{ iar Jacobianul este} \\ \frac{\partial g(\theta)}{\partial \theta} = \begin{pmatrix} \frac{\partial g(\theta)}{\partial \theta_1} & \frac{\partial g(\theta)}{\partial \theta_1} \end{pmatrix} = \begin{pmatrix} \frac{\partial g(\theta)}{\partial A} & \frac{\partial g(\theta)}{\partial \sigma^2} \end{pmatrix} = \begin{pmatrix} \frac{2A}{\sigma^2} & -\frac{A^2}{\sigma^4} \end{pmatrix}, \\ \frac{\partial g(\theta)}{\partial \theta} I^{-1}(\theta) \frac{\partial g(\theta)^t}{\partial \theta} = \begin{pmatrix} \frac{2A}{\sigma^2} & -\frac{A^2}{\sigma^4} \end{pmatrix} \begin{pmatrix} \frac{n}{\sigma^2} & 0 \\ 0 & \frac{n}{2\sigma^4} \end{pmatrix} \begin{pmatrix} \frac{2A}{\sigma^2} \\ -\frac{A^2}{\sigma^4} \end{pmatrix} = \frac{4A^2}{n\sigma^2} + \frac{2A^4}{n\sigma^4} = \frac{4\alpha + 2\alpha^2}{n}. \end{array} \right.$$

Cum α este scalar,

$$D^2(\hat{\alpha}) \geq \frac{4\alpha + 2\alpha^2}{n}.$$

Așa cum deja am discutat, eficiența estimatorului se păstrează în urma transformărilor liniare de tipul

$$\alpha = g(\theta) = \mathbf{A}\theta + \mathbf{b}, \quad \mathbf{A} \in \mathcal{M}_{r \times p}(\mathbb{R}), \quad \mathbf{b} \in \mathcal{M}_{r \times 1}(\mathbb{R}).$$

Dacă $\hat{\alpha} = \mathbf{A}\hat{\theta} + \mathbf{b}$, iar $\hat{\theta}$ este un estimator eficient (adică $C_{\hat{\theta}} = I^{-1}(\theta)$), atunci $\mathbb{E}(\hat{\alpha}) = \mathbf{A}\theta + \mathbf{b} = \alpha$ (adică $\hat{\alpha}$ este nedepășat) și avem:

$$C_{\hat{\alpha}} = \mathbf{A}C_{\hat{\theta}}\mathbf{A}^t = \mathbf{A}I^{-1}(\theta)\mathbf{A}^t = \frac{\partial g(\theta)}{\partial \theta} I^{-1}(\theta) \frac{\partial g(\theta)^t}{\partial \theta} = CRLB(\hat{\alpha}).$$

Pentru transformările neliniare, eficiența este menținută "aproximativ" doar pentru $n \rightarrow +\infty$. Aceasta înseamnă că, pentru un volum de selecție suficient de mare, densitatea de repartiție a estimatorului $\hat{\theta}$ este concentrată în jurul valorii reale, teoretice, a parametrului θ , adică estimatorul este consistent.

1.5 CRLB în cazul Gaussian general

În cazul în care observațiile sunt de tip Gaussian, putem deduce o formulă pentru CRLB care să generalizeze (9). Presupunem, pentru aceasta, că

$$V \sim \mathcal{N}(\mu(\theta), C(\theta)), \quad \text{unde } \mu(\theta) \in \mathcal{M}_{n \times 1}(\mathbb{R}) \quad \text{și} \quad C(\theta) \in \mathcal{M}_{n \times n}(\mathbb{R}).$$

Este evident că media și dispersia teoretice depind de parametrul θ .

Informația Fisher va fi:

$$(I(\theta))_{ij} = \left(\frac{\partial \mu(\theta)}{\partial \theta_i} \right)^t C^{-1}(\theta) \left(\frac{\partial \mu(\theta)}{\partial \theta_j} \right) + \frac{1}{2} \text{tr} \left(C^{-1}(\theta) \frac{\partial C(\theta)}{\partial \theta_i} C^{-1}(\theta) \frac{\partial C(\theta)}{\partial \theta_j} \right), \quad \text{unde} \quad (25)$$

$$\frac{\partial \mu(\theta)}{\partial \theta_i} = \begin{pmatrix} \frac{\partial (\mu(\theta))_1}{\partial \theta_i} \\ \frac{\partial (\mu(\theta))_2}{\partial \theta_i} \\ \vdots \\ \frac{\partial (\mu(\theta))_n}{\partial \theta_i} \end{pmatrix} \quad \text{și} \quad \frac{\partial C(\theta)}{\partial \theta_i} = \begin{pmatrix} \frac{\partial (C(\theta))_{11}}{\partial \theta_i} & \frac{\partial (C(\theta))_{12}}{\partial \theta_i} & \dots & \frac{\partial (C(\theta))_{1n}}{\partial \theta_i} \\ \frac{\partial (C(\theta))_{21}}{\partial \theta_i} & \frac{\partial (C(\theta))_{22}}{\partial \theta_i} & \dots & \frac{\partial (C(\theta))_{2n}}{\partial \theta_i} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{\partial (C(\theta))_{n1}}{\partial \theta_i} & \frac{\partial (C(\theta))_{n2}}{\partial \theta_i} & \dots & \frac{\partial (C(\theta))_{nn}}{\partial \theta_i} \end{pmatrix}.$$

Pentru situația $V \sim \mathcal{N}(\mu(\theta), C(\theta))$, aceasta se reduce la:

$$I(\theta) = \left(\frac{\partial \mu(\theta)}{\partial \theta} \right)^t C^{-1}(\theta) \left(\frac{\partial \mu(\theta)}{\partial \theta} \right) + \frac{1}{2} \text{tr} \left(\left(C^{-1}(\theta) \frac{\partial C(\theta)}{\partial \theta_i} \right)^2 \right), \quad (26)$$

după cum vedem în rezultatul tehnic ce urmează.

Densitatea de repartiție a vectorului aleator de selecție este:

$$f(v, \theta) = \frac{1}{(2\pi)^{n/2} \sqrt{\det(C(\theta))}} \exp \left(-\frac{1}{2} (v - \mu(\theta))^t C^{-1}(\theta) (v - \mu(\theta)) \right).$$

Vom utiliza următoarele identități:

$$\frac{\partial \ln(\det(C(\theta)))}{\partial \theta_k} = \text{tr} \left(C^{-1}(\theta) \frac{\partial C(\theta)}{\partial \theta_k} \right) \quad \text{și} \quad \frac{\partial C^{-1}(\theta)}{\partial \theta_k} = -C^{-1}(\theta) \frac{\partial C(\theta)}{\partial \theta_k} C^{-1}(\theta), \quad (27)$$

unde $\partial C(\theta)/\partial \theta_k$ este matricea de tipul $n \times n$, al cărei element de pe poziția (i, j) este $\partial(C(\theta))_{ij}/\partial \theta_k$. Pentru a obține prima formulă din (27), observăm, pentru început că

$$\frac{\partial \ln(\det(C(\theta)))}{\partial \theta_k} = \frac{1}{\det(C(\theta))} \frac{\partial \det(C(\theta))}{\partial \theta_k}. \quad (28)$$

Cum $\det(C(\theta))$ depinde de toate elementele matricei $C(\theta)$,

$$\frac{\partial \det(C(\theta))}{\partial \theta_k} = \sum_{i=1}^n \sum_{j=1}^n \frac{\partial \det(C(\theta))}{\partial(C(\theta))_{ij}} \frac{\partial(C(\theta))_{ij}}{\partial \theta_k} = \text{tr} \left(\frac{\partial \det(C(\theta))}{\partial C(\theta)} \frac{\partial C(\theta)}{\partial \theta_k} \right), \quad (29)$$

unde $\partial \det(C(\theta))/\partial C(\theta)$ este o matrice de tip $n \times n$, al cărei element de pe poziția (i, j) este $\partial \det(C(\theta))/\partial(C(\theta))_{ij}$. Am ținut cont și de identitatea matriceală $\text{tr}(AB^t) = \sum_{i=1}^n \sum_{j=1}^n A_{ij} B_{ij}$. Din definiția determinantului,

$$\det(C(\theta)) = \sum_{i=1}^n (C(\theta))_{ij} M_{ij}, \quad M \in \mathcal{M}_{n \times n}(\mathbb{R}) \text{ este matricea cofactor și } j \text{ este arbitrar în } \{1, 2, \dots, n\}.$$

Prin urmare,

$$\frac{\partial \det(C(\theta))}{\partial(C(\theta))_{ij}} = M_{ij}, \quad \text{sau, echivalent,} \quad \frac{\partial \det(C(\theta))}{\partial C(\theta)} = M.$$

Pe de altă parte,

$$C^{-1}(\theta) = \frac{M^t}{\det(C(\theta))} \quad \text{și deci} \quad \frac{\partial \det(C(\theta))}{\partial C(\theta)} = C^{-1}(\theta) \det(C(\theta))$$

Folosim acum (28) și (29) și obținem rezultatul dorit:

$$\frac{\partial \ln(\det(C(\theta)))}{\partial \theta_k} = \frac{1}{\det(C(\theta))} \text{tr} \left(C^{-1}(\theta) \det(C(\theta)) \frac{\partial C(\theta)}{\partial \theta_k} \right) = \text{tr} \left(C^{-1}(\theta) \frac{\partial C(\theta)}{\partial \theta_k} \right).$$

Pentru cea de a doua identitate din (27), plecăm de la egalitatea evidentă $C^{-1}(\theta)C(\theta) = I$. Diferențiind componentele matricelor și scriind sub formă matriceală, obținem:

$$C^{-1}(\theta) \frac{\partial C(\theta)}{\partial \theta_k} + \frac{\partial C^{-1}(\theta)}{\partial \theta_k} C(\theta) = \mathbf{0},$$

iar concluzia rezultă imediat.

Ne preocupăm acum de determinarea CRLB. Derivatele parțiale în raport cu parametrii conduc la următoarea formă a funcției de verosimilitate:

$$\frac{\partial \ln f(V, \theta)}{\partial \theta_k} = -\frac{1}{2} \frac{\partial \ln \det(C(\theta))}{\partial \theta_k} - \frac{1}{2} \frac{\partial}{\partial \theta_k} \left((V - \mu(\theta))^t C^{-1}(\theta) (V - \mu(\theta)) \right).$$

Cum primul termen a fost evaluat în (27), ne vom ocupa de cel de al doilea:

$$\begin{aligned} \frac{\partial}{\partial \theta_k} \left((V - \mu(\theta))^t C^{-1}(\theta) (V - \mu(\theta)) \right) &= \frac{\partial}{\partial \theta_k} \sum_{i=1}^n \sum_{j=1}^n \left((X_i - (\mu(\theta))_i) (C^{-1}(\theta))_{ij} (X_j - (\mu(\theta))_j) \right) \\ &= \sum_{i=1}^n \sum_{j=1}^n \left\{ (X_i - (\mu(\theta))_i) \left[(C^{-1}(\theta))_{ij} \left(-\frac{\partial (\mu(\theta))_j}{\partial \theta_k} \right) + \frac{\partial (C^{-1}(\theta))_{ij}}{\partial \theta_k} (X_j - (\mu(\theta))_j) \right] \right. \\ &\quad \left. + \left(-\frac{\partial (\mu(\theta))_i}{\partial \theta_k} \right) (C^{-1}(\theta))_{ij} (X_j - (\mu(\theta))_j) \right\} \\ &= -(V - \mu(\theta))^t C^{-1}(\theta) \frac{\partial \mu(\theta)}{\partial \theta_k} + (V - \mu(\theta))^t \frac{\partial C^{-1}(\theta)}{\partial \theta_k} (V - \mu(\theta)) - \frac{\partial \mu(\theta)^t}{\partial \theta_k} C^{-1}(\theta) (V - \mu(\theta)) \\ &= -2 \frac{\partial \mu(\theta)^t}{\partial \theta_k} C^{-1}(\theta) (V - \mu(\theta)) + (V - \mu(\theta))^t \frac{\partial C^{-1}(\theta)}{\partial \theta_k} (V - \mu(\theta)). \end{aligned}$$

Folosind (27) împreună cu precedentul rezultat obținem:

$$\frac{\partial \ln f(V, \theta)}{\partial \theta_k} = -\frac{1}{2} \text{tr} \left(C^{-1}(\theta) \frac{\partial C(\theta)}{\partial \theta_k} \right) + \frac{\partial \mu(\theta)^t}{\partial \theta_k} C^{-1}(\theta) (V - \mu(\theta)) - \frac{1}{2} (V - \mu(\theta))^t \frac{\partial C^{-1}(\theta)}{\partial \theta_k} (V - \mu(\theta)). \quad (30)$$

Notăm acum $Y := V - \mu(\theta)$. Pentru a determina CRLB trebuie să evaluăm elementele

$$(I(\theta))_{kl} = \mathbb{E} \left(\frac{\partial \ln f(V, \theta)}{\partial \theta_k} \frac{\partial \ln f(V, \theta)}{\partial \theta_l} \right), \quad k, l \in \{1, 2, \dots, p\}.$$

Avem, astfel,

$$\begin{aligned} (I(\theta))_{kl} &= \frac{1}{4} \text{tr} \left(C^{-1}(\theta) \frac{\partial C(\theta)}{\partial \theta_k} \right) \text{tr} \left(C^{-1}(\theta) \frac{\partial C(\theta)}{\partial \theta_l} \right) + \frac{1}{2} \text{tr} \left(C^{-1}(\theta) \frac{\partial C(\theta)}{\partial \theta_k} \right) \mathbb{E} \left(Y^t \frac{\partial C^{-1}(\theta)}{\partial \theta_l} Y \right) \\ &\quad + \frac{\partial \mu(\theta)^t}{\partial \theta_k} C^{-1}(\theta) \mathbb{E} (Y Y^t) C^{-1}(\theta) \frac{\partial \mu(\theta)}{\partial \theta_l} + \frac{1}{4} \mathbb{E} \left(Y^t \frac{\partial C^{-1}(\theta)}{\partial \theta_k} Y Y^t \frac{\partial C^{-1}(\theta)}{\partial \theta_l} Y \right), \end{aligned}$$

toate momentele de ordin impar fiind zero. Continuând evaluarea,

$$\begin{aligned} (I(\theta))_{kl} &= \frac{1}{4} \text{tr} \left(C^{-1}(\theta) \frac{\partial C(\theta)}{\partial \theta_k} \right) \text{tr} \left(C^{-1}(\theta) \frac{\partial C(\theta)}{\partial \theta_l} \right) - \frac{1}{2} \text{tr} \left(C^{-1}(\theta) \frac{\partial C(\theta)}{\partial \theta_k} \right) \text{tr} \left(C^{-1}(\theta) \frac{\partial C(\theta)}{\partial \theta_l} \right) \\ &\quad + \frac{\partial \mu(\theta)^t}{\partial \theta_k} C^{-1}(\theta) \frac{\partial \mu(\theta)}{\partial \theta_l} + \frac{1}{4} \mathbb{E} \left(Y^t \frac{\partial C^{-1}(\theta)}{\partial \theta_k} Y Y^t \frac{\partial C^{-1}(\theta)}{\partial \theta_l} Y \right). \end{aligned} \quad (31)$$

Am utilizat (27) și faptul că $\mathbb{E}(Y^t Z) = \text{tr}(\mathbb{E}(Z Y^t))$, pentru $Y, Z \in \mathcal{M}_{n \times 1}(\mathbb{R})$. Pentru estimarea ultimului termen, folosim un rezultat furnizat de Porat, Friedlander, 1986:

$$\mathbb{E}(Y^t A Y Y^t B Y) = \text{tr}(A C) \text{tr}(B C) + 2 \text{tr}(A C B C), \quad \text{unde } C := \mathbb{E}(Y Y^t),$$

iar matricele A și B sunt simetrice. În virtutea relației anterioare și, folosind din nou relația (27), ultimul termen din reprezentarea (31) devine:

$$\begin{aligned} &\frac{1}{4} \text{tr} \left(\frac{\partial C^{-1}(\theta)}{\partial \theta_k} C(\theta) \right) \text{tr} \left(\frac{\partial C^{-1}(\theta)}{\partial \theta_l} C(\theta) \right) + \frac{1}{2} \text{tr} \left(\frac{\partial C^{-1}(\theta)}{\partial \theta_k} C(\theta) \frac{\partial C^{-1}(\theta)}{\partial \theta_l} C(\theta) \right) \\ &= \frac{1}{4} \text{tr} \left(C^{-1}(\theta) \frac{\partial C(\theta)}{\partial \theta_k} \right) \text{tr} \left(C^{-1}(\theta) \frac{\partial C(\theta)}{\partial \theta_l} \right) + \frac{1}{2} \text{tr} \left(C^{-1}(\theta) \frac{\partial C(\theta)}{\partial \theta_k} C^{-1}(\theta) \frac{\partial C(\theta)}{\partial \theta_l} \right). \end{aligned} \quad (32)$$

Inserăm (32) în (31) și obținem rezultatul dorit.

Exemplul 1.7 Parametrii unui semnal în WGN.

I. Presupunem că dorim să estimăm un parametru scalar σ^2 pentru setul de date

$$X_i = s_i(\theta) + W_i, \quad \text{unde } W_i \text{ este WGN, } i \in \{0, 1, \dots, n-1\}.$$

Matricea de covarianță este $C = -I$ și este independentă de parametrul θ , ceea ce face ca al doilea termen din (26) să fie zero. Obținem astfel:

$$I(\theta) = \frac{1}{\sigma^2} \left(\left(\frac{\partial \mu(\theta)}{\partial \theta} \right)^t \left(\frac{\partial \mu(\theta)}{\partial \theta} \right) \right) = \frac{1}{\sigma^2} \sum_{i=0}^{n-1} \left(\frac{\partial (\mu(\theta))_i}{\partial \theta} \right)^2 = \frac{1}{\sigma^2} \sum_{i=0}^{n-1} \left(\frac{\partial s_i(\theta)}{\partial \theta} \right)^2,$$

ceea ce este în concordanță cu (9). Dacă dorim să generalizăm la cazul unui parametru vectorial al semnalului, din (25) obținem:

$$(I(\theta))_{ij} = \frac{1}{\sigma^2} \sum_{k=0}^{n-1} \frac{\partial s_k(\theta)}{\partial \theta_i} \frac{\partial s_k(\theta)}{\partial \theta_j}.$$

II. Dacă observăm $X_i = W_i$, unde W_i este WGN cu dispersia $\theta = \sigma^2$, $i \in \{0, 1, \dots, n-1\}$, atunci, în conformitate cu (26), deoarece $C(\sigma^2) = \sigma^2 I$, avem:

$$I(\theta) = \frac{1}{2} \text{tr} \left(\left(C^{-1}(\sigma^2) \frac{\partial C(\sigma^2)}{\partial \sigma^2} \right)^2 \right) = \frac{1}{2} \text{tr} \left(\left(\frac{1}{\sigma^2} I \right)^2 \right) = \frac{1}{2\sigma^4} \text{tr}(I) = \frac{n}{2\sigma^4},$$

ceea ce este în concordanță cu Exemplul 1.5.

III. Presupunem datele:

$$X_i = A + W_i, \quad \text{unde } W_i \text{ este WGN, } i \in \{0, 1, \dots, n-1\},$$

iar A , mărimea (amplitudinea) curentului continuu, este o variabilă aleatoare Gaussiană, de medie 0 și dispersie σ_A^2 ($A \sim \mathcal{N}(0, \sigma_A^2)$), independentă de $i \in \{0, 1, \dots, n-1\}$. Puterea semnalului (dispersia σ_A^2) se presupune a fi parametrul necunoscut. Vectorul de selecție $V = (X_0, X_1, \dots, X_{n-1})^t$ va fi, prin urmare, un vector repartizat tot normal, de medie 0 și având matricea de covarianță de tipul $n \times n$, ale cărei elemente sunt:

$$(C(\sigma_A^2))_{ij} = \mathbb{E}(X_i X_j) = \mathbb{E}((A + W_i)(A + W_j)) = \sigma_A^2 + \sigma^2 \delta_{ij}.$$

Aceasta înseamnă, sub formă vectorială,

$$C(\sigma_A^2) = \sigma_A^2 \mathbf{1} \cdot \mathbf{1}^t + \sigma^2 I, \quad \text{unde } \mathbf{1} = (1, 1, \dots, 1)^t \in \mathcal{M}_{n \times 1}(\mathbb{R}).$$

Aplicăm acum identitatea Woodbury privind inversa unei matrice având o anumită structură particulară, dată (vezi Kay [2, Annex A1.1.3, Matrix Manipulation and Formulas]): dacă $A \in \mathcal{M}_{n \times n}(\mathbb{R})$ și $u \in \mathcal{M}_{n \times 1}(\mathbb{R})$, atunci are loc relația:

$$(A + uu^t)^{-1} = A^{-1} - \frac{A^{-1}uu^tA^{-1}}{1 + u^tA^{-1}u}.$$

Obținem astfel:

$$C^{-1}(\sigma_A^2) = \frac{1}{\sigma^2} \left(I - \frac{\sigma_A^2}{\sigma^2 + n\sigma_A^2} \mathbf{1} \cdot \mathbf{1}^t \right).$$

Deoarece

$$\frac{\partial C(\sigma_A^2)}{\partial \sigma_A^2} = \mathbf{1} \cdot \mathbf{1}^t, \quad \text{deducem } C^{-1}(\sigma_A^2) \frac{\partial C(\sigma_A^2)}{\partial \sigma_A^2} = \frac{1}{\sigma^2 + n\sigma_A^2} \mathbf{1} \cdot \mathbf{1}^t.$$

Înlocuim în (26) și avem:

$$I(\sigma_A^2) = \frac{1}{2} \text{tr} \left(\left(\frac{1}{\sigma^2 + n\sigma_A^2} \right)^2 \mathbf{1} \cdot \mathbf{1}^t \cdot \mathbf{1} \cdot \mathbf{1}^t \right) = \frac{n}{2} \left(\frac{1}{\sigma^2 + n\sigma_A^2} \right)^2 \text{tr}(\mathbf{1} \cdot \mathbf{1}^t) = \frac{1}{2} \left(\frac{n}{\sigma^2 + n\sigma_A^2} \right)^2,$$

adică CRLB este

$$D^2(\widehat{\sigma_A^2}) \geq 2 \left(\sigma_A^2 + \frac{\sigma^2}{n} \right)^2 = \text{CRLB}(\widehat{\sigma_A^2}).$$

Observăm că, chiar și pentru $n \rightarrow +\infty$, $\text{CRLB}(\widehat{\sigma_A^2})$ obținut nu poate descrește sub valoarea $2\sigma_A^4$. Aceasta se întâmplă datorită faptului că fiecare dată suplimentară în eșantion vine cu aceeași valoare pentru A .

1.6 CRLB asimptotică pentru procese aleatoare Gaussiene staționare în sens larg

Uneori, în practică, poate fi mai dificil de determinat CRLB folosind (25) datorită necesității determinării inversei matricei de covarianță. Deși aceasta se poate determina numeric, folosind un software matematic, există și o abordare alternativă, dar care poate fi aplicată doar proceselor Gaussiene staționare în sens larg (WSS, Wide Sense Stationary). Pentru aceasta, volumul de selecție $n \rightarrow +\infty$. Pentru procese cu densitatea spectrală de putere (PSD, Power Spectral Density) largă, aproximarea va fi bună pentru înregistrări de date de lungime moderată, în timp ce procese cu bandă îngustă trebuiesc înregistrări de date de lungime sporită.

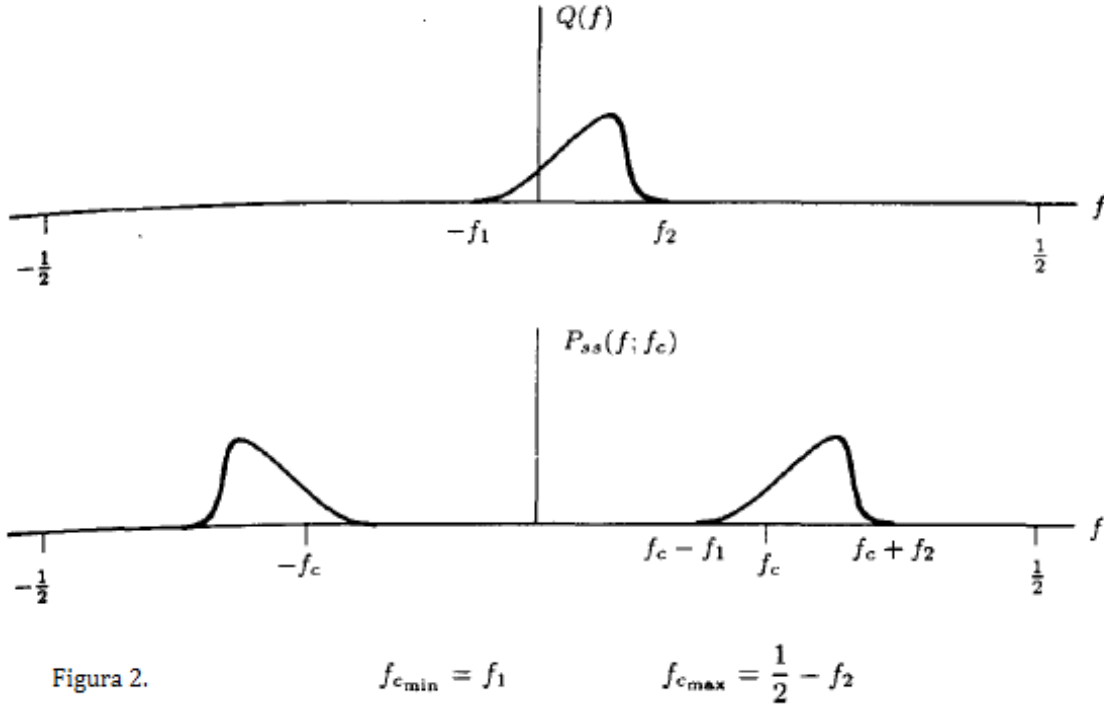
Elementele matricei informației Fisher sunt, asimptotic, pentru $n \rightarrow +\infty$,

$$(I(\theta))_{ij} = \frac{n}{2} \int_{-1/2}^{1/2} \frac{\partial \ln P_{xx}(f, \theta)}{\partial \theta_i} \frac{\partial \ln P_{xx}(f, \theta)}{\partial \theta_j} df, \quad (33)$$

unde $P_{xx}(f, \theta)$ este indicatorul PSD al procesului, cu dependența explicită de parametrul vectorial θ (vezi Kay [2, Appendix 3D]). Presupunem că $\mathbb{E}(X_i) = 0$, pentru fiecare i .

Exemplul 1.8 Frecvența centrală a procesului. O problemă uzuală constă în estimarea frecvenței centrale f_c a unei PSD, care, de altfel, este cunoscută. Presupunem că este dată:

$$P_{xx}(f, f_c) = Q(f - f_c) + Q(-f - f_c) + \sigma^2$$



și dorim să determinăm CRLB pentru f_c , în ipoteza că $Q(f)$ și σ^2 sunt cunoscute. Putem percepe procesul ca fiind un semnal aleator aflat într-un WGN. Frecvențele centrale posibile sunt constrânse a fi situate în intervalul $[f_1, 1/2 - f_2]$. Pentru aceste frecvențe centrale, PSD-ul semnalului, pentru $f \geq 0$, se va situa în intervalul $[0, 1/2]$. Atunci, cum $\theta = f_c$ este un scalar, din (33), obținem:

$$D^2(\hat{f}_c) \geq \frac{1}{\frac{n}{2} \int_{-1/2}^{1/2} \left(\frac{\partial \ln P_{xx}(f, f_c)}{\partial f_c} \right)^2 df}.$$

Dar,

$$\frac{\partial \ln P_{xx}(f, f_c)}{\partial f_c} = \frac{\partial \ln (Q(f - f_c) + Q(-f - f_c) + \sigma^2)}{\partial f_c} = \frac{\frac{\partial Q(f - f_c)}{\partial f_c} + \frac{\partial Q(-f - f_c)}{\partial f_c}}{Q(f - f_c) + Q(-f - f_c) + \sigma^2},$$

care, fiind o funcție impară în raport cu f , conduce la:

$$\int_{-1/2}^{1/2} \left(\frac{\partial \ln P_{xx}(f, f_c)}{\partial f_c} \right)^2 df = 2 \int_0^{1/2} \left(\frac{\partial \ln P_{xx}(f, f_c)}{\partial f_c} \right)^2 df.$$

De asemenea, pentru $f \geq 0$, $Q(-f - f_c) = 0$, iar derivata sa va fi egală cu zero în virtutea celor discutate privind restricțiile. Rezultă că:

$$\begin{aligned} D^2(\hat{f}_c) &\geq \frac{1}{n \int_0^{1/2} \left(\frac{\partial Q(f - f_c) / \partial f_c}{Q(f - f_c) + \sigma^2} \right)^2 df} = \frac{1}{n \int_0^{1/2} \left(\frac{-\partial Q(f - f_c) / \partial (f - f_c)}{Q(f - f_c) + \sigma^2} \right)^2 df} \\ &= \frac{1}{n \int_{-f_c}^{1/2 - f_c} \left(\frac{-\partial Q(f') / \partial f'}{Q(f') + \sigma^2} \right)^2 df'}, \quad \text{unde } f' = f - f_c. \end{aligned}$$

Dar $1/2 - f_c \geq 1/2 - f_{c\max} = f_2$ și $-f_c \leq -f_{c\min} = -f_1$, putând astfel schimba capetele de integrare în intervalul $[-1/2, 1/2]$. Obținem:

$$D^2(\hat{f}_c) \geq \frac{1}{n \int_{-1/2}^{1/2} \left(\frac{\partial Q(f) / \partial f}{Q(f) + \sigma^2} \right)^2 df} = \frac{1}{n \int_{-1/2}^{1/2} \left(\frac{\partial \ln (Q(f) + \sigma^2)}{\partial f} \right)^2 df}.$$

Dacă considerăm, ca un caz particular,

$$Q(f) := \exp\left(-\frac{1}{2}\left(\frac{f}{\sigma_f}\right)^2\right), \quad \text{unde } \sigma_f \ll 1/2,$$

atunci $Q(f) \gg \sigma^2$ și avem, asimptotic,

$$D^2(\hat{f}_c) \geq \frac{1}{n \int_{-1/2}^{1/2} \frac{f^2}{\sigma_f^4} df} = \frac{12\sigma_f^4}{n}.$$

Lățimi de bandă mai mici decât σ_f^4 vor conduce la limite inferioare de mărginire mai joase pentru frecvența centrală deoarece PSD-ul se modifică mai rapid odată cu schimbările lui f_c .

1.6.1 Exemple privind procesarea semnalelor

Aplicăm teoria deducerii CRLB pentru câteva probleme de interes privind procesarea semnalelor: *estimarea ariei de acoperire, estimarea poziției țintei (sonar, radar, robotică), estimarea frecvenței (sonar, radar, econometrie, spectrometrie), estimarea parametrilor autoregresivi (teoria limbajului, econometrie).*

Exemplul 1.9 Estimarea range-ului (raza de acțiune). Un radar emite un semnal de tip puls. Timpul total al traseului semnalului de la emițător la țintă și înapoi va fi notat cu τ_0 și este dependent de range (R) prin formula $\tau_0 = 2R/c$, unde c este viteza sa de propagare. Estimarea range-ului este deci echivalentă cu estimarea timpului de întoarcere a semnalului la sursa emitentă. Dacă $s(t)$ este semnalul transmis, un model simplu pentru unda continuă recepționată este:

$$x(t) = s(t - \tau_0) + W(t), \quad t \in [0, T].$$

Presupunem că pulsul transmis este diferit de zero pe $[0, T_s]$ și este limitat la o lățime de bandă de B Hz. Dacă întârzierea de timp până la recepționarea semnalului revenit este $\tau_{0,\max}$, atunci intervalul de observare trebuie ales astfel încât să includă întregul semnal, adică $T = T_s + \tau_{0,\max}$. Zgomotul pe sistem este de tip Gaussian, cu PSD-ul și autocorelația (ACF-ul - corelația unui semnal cu o copie întârziată a sa, ca funcție de întârziere. Informal, este similaritatea dintre observații, în funcție de intervalul de timp dintre ele) ca în Figura 3.

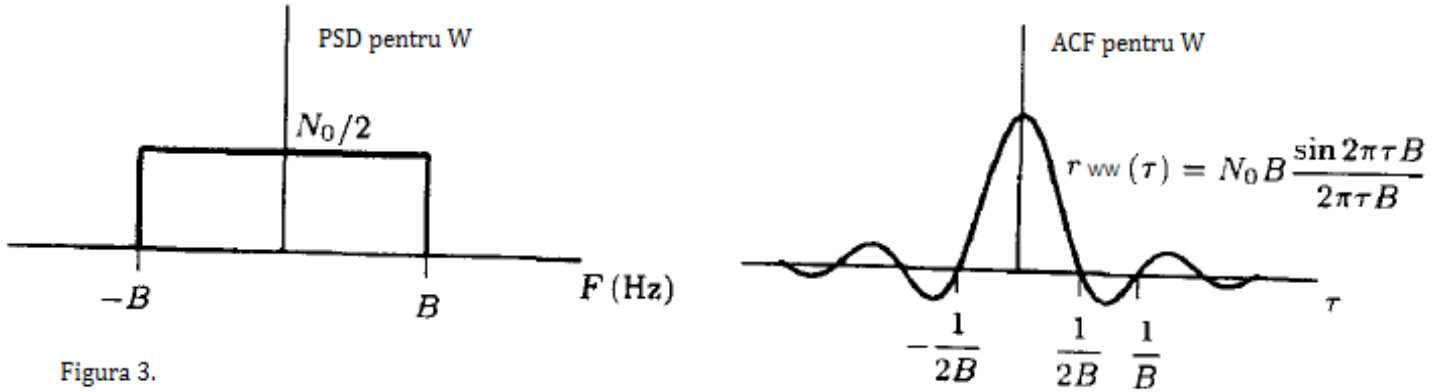


Figura 3.

Limitarea de bandă a zgomotului rezultă prin filtrarea undei la lățimea de bandă a semnalului de B Hz. Pentru eșantionare, considerăm $\Delta := 1/(2B)$, iar datele observate vor fi:

$$X_{i\Delta} = s(i\Delta - \tau_0) + W(i\Delta), \quad i \in \{0, 1, \dots, n-1\}.$$

Notăm cu X_i și W_i elementele selecției și construim șirul de date discret:

$$X_i = s(i\Delta - \tau_0) + W_i. \quad (34)$$

Menționăm că, pentru fiecare i , W_i sunt WGN deoarece observațiile sunt separate de $k\Delta = k/(2B)$, care corespunde zerourilor funcției ACF a lui $W(t)$, după cum se observă în Figura 3. De asemenea,

$$D^2(W_i) = \sigma^2(= r_{WW}(0)) = N_0 B.$$

Deoarece semnalul este nenul doar pe intervalul $\tau_0 \leq t \leq \tau_0 + T_s$, modelul (34) se reduce la:

$$X_i = \begin{cases} W_i, & 1 \leq i \leq i_0 - 1, \\ s(i\Delta - \tau_0) + W_i, & i_0 \leq i \leq i_0 + M - 1, \\ W_i, & i_0 + M \leq i \leq n, \end{cases} \quad (35)$$

unde M este lungimea semnalului eșantionat, iar $i_0 = \tau_0/\Delta$ reprezintă întârzierea, în eșantioane. Pentru simplitate, se presupune Δ mic astfel încât τ_0/Δ să poată fi aproximat cu un întreg. Cu această formulare, putem aplica (9) pentru evaluarea CRLB:

$$D^2(\hat{\tau}_0) \geq \frac{\sigma^2}{\sum_{i=0}^{n-1} \left(\frac{\partial s(i, \tau_0)}{\partial \tau_0} \right)^2} = \frac{\sigma^2}{\sum_{i=i_0}^{i_0+M-1} \left(\frac{\partial s(i\Delta - \tau_0)}{\partial \tau_0} \right)^2} = \frac{\sigma^2}{\sum_{i=i_0}^{i_0+M-1} \left(\frac{ds(t)}{dt} \Big|_{t=i\Delta-\tau_0} \right)^2} = \frac{\sigma^2}{\sum_{i=0}^{M-1} \left(\frac{ds(t)}{dt} \Big|_{t=i\Delta} \right)^2},$$

deoarece $\tau_0 = i_0\Delta$. Cum Δ este suficient de mic, aproximăm suma cu o integrală și obținem:

$$D^2(\hat{\tau}_0) \geq \frac{\sigma^2}{\frac{1}{\Delta} \int_0^{T_s} \left(\frac{ds(t)}{dt} \right)^2 dt} = \frac{N_0/2}{\int_0^{T_s} (s'(t))^2 dt}, \quad \text{deoarece } \Delta = 1/(2B) \text{ și } \sigma^2 = N_0B. \quad (36)$$

Cum energia poate fi reprezentată integral, $\mathcal{E} = \int_0^{T_s} s^2(t) dt$, inegalitatea (36) se poate scrie, sub formă echivalentă,

$$D^2(\hat{\tau}_0) \geq \frac{1}{\frac{\mathcal{E}}{N_0/2} \overline{F^2}}, \quad \text{unde } \overline{F^2} = \frac{\int_0^{T_s} (s'(t))^2 dt}{\int_0^{T_s} s^2(t) dt}.$$

$\overline{F^2}$ este o măsură a lățimii de bandă a semnalului. Cu cât este mai mare această cantitate, cu atât mai mică va fi CRLB.

Exemplul 1.10 Estimarea parametrului vectorial al semnalului sinusoidal. În multe aplicații se pune problema estimării parametrilor unui semnal de tip sinusoidal. Sonarele și mecanismele de tip radar pot cataloga semnalul observat ca fiind unul de tip sinusoidal. Ne punem problema determinării CRLB pentru amplitudinea A , frecvența f_0 și faza ϕ a unui curent sinusoidal într-un WGN. Datele sunt de tipul:

$$X_i = A \cos(2\pi f_0 i + \phi) + W_i, \quad i = 0, 1, \dots, n-1,$$

unde $A > 0$ și $0 < f_0 < 1/2$. Deoarece avem mai mulți parametri necunoscuți, folosim formula:

$$(I(\theta))_{ij} = \frac{1}{\sigma^2} \sum_{k=0}^{n-1} \frac{\partial s_k(\theta)}{\partial \theta_i} \frac{\partial s_k(\theta)}{\partial \theta_j}, \quad \text{pentru } \theta = (A, f_0, \phi)^t.$$

Pentru evaluarea CRLB, presupunem că f_0 nu este în vecinătatea lui 0 sau 1/2, fapt ce ne permite să utilizăm aproximările următoare (vezi Stoica, P.; R.L. Moses; B. Friedlander; T. Soderstrom, Maximum Likelihood Estimation of the Parameters of Multiple Sinusoids from Noisy Measurements, IEEE Trans. Acoust., Speech, Signal Process., Vol. 37, pp. 378-392, March 1989):

$$\frac{1}{n^{i+1}} \sum_{k=0}^{n-1} k^i \sin(4\pi f_0 k + 2\phi) \approx 0 \quad \text{și} \quad \frac{1}{n^{i+1}} \sum_{k=0}^{n-1} k^i \cos(4\pi f_0 k + 2\phi) \approx 0, \quad i \in \{0, 1, 2\}.$$

Folosind aceste aproximații și notând $\alpha := 2\pi f_0 i + \phi$, avem:

$$\left\{ \begin{array}{l} (I(\theta))_{11} = \frac{1}{\sigma^2} \sum_{k=0}^{n-1} \cos^2 \alpha = \frac{1}{\sigma^2} \sum_{k=0}^{n-1} \left(\frac{1}{2} + \frac{1}{2} \cos 2\alpha \right) \approx \frac{n}{2\sigma^2}, \\ (I(\theta))_{12} = -\frac{1}{\sigma^2} \sum_{k=0}^{n-1} 2A\pi k \cos \alpha \sin \alpha = -\frac{\pi A}{\sigma^2} \sum_{k=0}^{n-1} k \sin 2\alpha \approx 0, \\ (I(\theta))_{13} = -\frac{1}{\sigma^2} \sum_{k=0}^{n-1} A \cos \alpha \sin \alpha = -\frac{A}{2\sigma^2} \sum_{k=0}^{n-1} \sin 2\alpha \approx 0, \\ (I(\theta))_{22} = \frac{1}{\sigma^2} \sum_{k=0}^{n-1} A^2 (2\pi k)^2 \sin^2 \alpha = \frac{(2\pi A)^2}{\sigma^2} \sum_{k=0}^{n-1} k^2 \left(\frac{1}{2} - \frac{1}{2} \cos 2\alpha \right) \approx \frac{(2\pi A)^2}{2\sigma^2} \sum_{k=0}^{n-1} k^2, \\ (I(\theta))_{23} = \frac{1}{\sigma^2} \sum_{k=0}^{n-1} A^2 2\pi k \sin^2 \alpha \approx \frac{\pi A^2}{\sigma^2} \sum_{k=0}^{n-1} k, \\ (I(\theta))_{33} = \frac{1}{\sigma^2} \sum_{k=0}^{n-1} A^2 \sin^2 \alpha \approx \frac{nA^2}{2\sigma^2}. \end{array} \right.$$

Matricea informației Fisher este:

$$I(\theta) = \frac{1}{\sigma^2} \begin{pmatrix} \frac{n}{2} & 0 & 0 \\ 0 & 2\pi^2 A^2 \sum_{k=0}^{n-1} k^2 & \pi A^2 \sum_{k=0}^{n-1} k \\ 0 & \pi A^2 \sum_{k=0}^{n-1} k & \frac{nA^2}{2} \end{pmatrix}.$$

Obținem astfel marginile inferioare:

$$D^2(\hat{A}) \geq \frac{2\sigma^2}{n}, \quad D^2(\hat{f}_0) \geq \frac{12}{(2\pi)^2 \eta n (n^2 - 1)} \quad \text{și} \quad D^2(\hat{\phi}) \geq \frac{2(2n-1)}{\eta n (n+1)}, \quad (37)$$

unde $\eta := A^2 / (2\sigma^2)$ este coeficientul SNR (signal-to-noise ratio) al semnalului.

Exemplul 1.11 Estimarea menținerii țintei semnalului. În problemele de localizare, este importantă estimarea modului în care este identificată poziția unei ținte. Presupunem câmpul acustic de presiune fiind format dintr-un șir de senzori plasați echidistant și că ținta radiază un semnal de tip sinusoidal de forma $A \cos(2\pi F_0 t + \phi)$. Semnalul recepționat de senzorul cu indicele i este $A \cos(2\pi F_0 (t - t_i) + \phi)$, unde t_i este timpul de propagare până la senzorul indicele i . Așa cum se observă în Figura 4, distanța temporală a frontul de undă la senzorul marcat cu $i - 1$ diferă față de cea de la senzorul marcat cu i prin $d \cos \beta / c$, datorită distanței suplimentare de propagare.

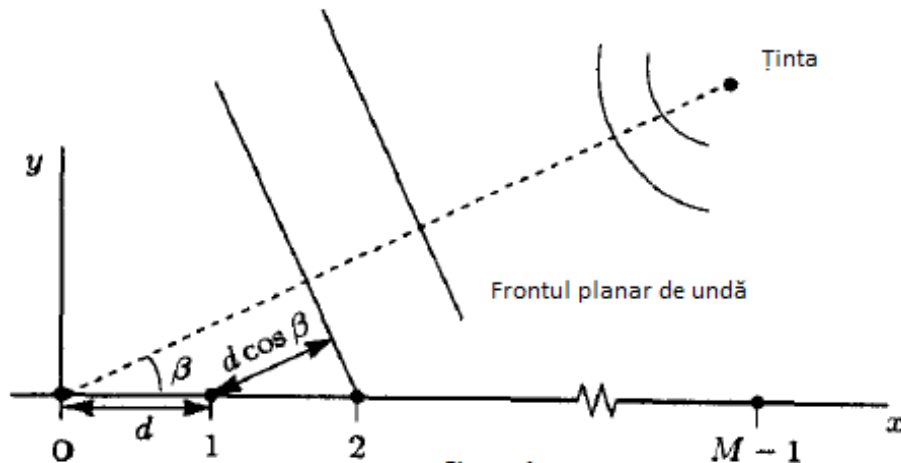


Figura 4.

Prin urmare, timpul de propagare până la senzorul i este:

$$t_i = t_0 - i \frac{d}{c} \cos \beta, \quad i \in \{0, 1, \dots, M-1\},$$

unde t_0 este timpul de propagare până la senzorul cu indicele 0. Semnalul observat la senzorul cu indicele i este

$$s_i(t) = A \cos \left(2\pi F_0 \left(t - t_0 + i \frac{d}{c} \cos \beta \right) + \phi \right).$$

Dacă se realizează o singură "fotografie" a datelor, adică ne interesează doar informația de la un moment fixat t_s ,

$$s_i(t_s) = A \cos \left(2\pi \left(F_0 \frac{d}{c} \cos \beta \right) i + \phi' \right), \quad \text{unde } \phi' = \phi + 2\pi F_0 (t_s - t_0). \quad (38)$$

Sub această formă, este clar că observațiile spațiale sunt de tip sinusoidal, cu frecvența $f_s = F_0 (d/c) \cos \beta$. Presupunem, de asemenea, că datele furnizate de senzor sunt perturbate de către un zgomot Gaussian de medie 0 și dispersie σ^2 , zgomotele ce sunt independente stochastic de la un senzor la altul. Datele eșantionului vor fi modelate prin variabilele de selecție:

$$X_i = s_i(t_s) + W_i, \quad W_i \text{ este WGN, pentru } i \in \{0, 1, 2, \dots, M-1\}.$$

Deoarece, de regulă, A, ϕ, β sunt parametri necunoscuți, se pune problema estimării parametrilor $\{A, f_s, \phi'\}$ din (38), ca în Exemplul 1.10. Odată ce am determinat CRLB pentru aceștia, putem folosi formulele pentru transformarea parametrilor. Transformarea este $\theta = (A, f_s, \phi')^t$:

$$\alpha = g(\theta) = \begin{pmatrix} A \\ \beta \\ \phi' \end{pmatrix} = \begin{pmatrix} A \\ \arccos \left(\frac{cf_s}{F_0 d} \right) \\ \phi' \end{pmatrix}, \quad \text{cu Jacobianul } \frac{\partial g(\theta)}{\partial \theta} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & -\frac{c}{F_0 d \sin \beta} & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

În conformitate cu (24), avem:

$$\left(C_{\hat{\alpha}} - \frac{\partial g(\theta)}{\partial \theta} I^{-1}(\theta) \frac{\partial g(\theta)^t}{\partial \theta} \right)_{22} \geq 0.$$

Forma diagonală a Jacobianului implică:

$$D^2(\hat{\beta}) \geq \left(\frac{\partial g(\theta)}{\partial \theta} \right)_{22}^2 (I^{-1}(\theta))_{22}.$$

Formula (37) afirmă că $(I^{-1}(\theta))_{22} = 12 / \left[(2\pi)^2 \eta M (M^2 - 1) \right]$ și, prin urmare,

$$D^2(\hat{\beta}) \geq \frac{12}{(2\pi)^2 \eta M (M^2 - 1)} \frac{c^2}{F_0^2 d^2 \sin^2 \beta} = \frac{12}{(2\pi)^2 \eta M \frac{M+1}{M-1} \left(\frac{L}{\lambda} \right)^2 \sin^2 \beta}, \quad (39)$$

unde $\lambda = c/F_0$ este lungimea de undă a valului propagant planar și $L = (M-1)d$ este lungimea șirului de senzori. Estimarea este cel mai ușor de realizat dacă $\beta = \pi/2$ și imposibil de realizat dacă $\beta = 0$.

2 Estimatori nedeplasați de varianță minimă (estimatori MVU)

2.1 Introducere

Am observat că, în evaluarea CRLB, obținem uneori un estimator eficient și, în consecință, acesta este nedeplasat și de dispersie (varianță) minimă. Dacă, în schimb, nu există un estimator eficient, se poate formula totuși problema existenței și identificării unui estimator MVU. Pentru a realiza acest obiectiv, introducem conceptul de statistici suficiente și, ca instrument de lucru, Teorema Rao-Blackwell-Lehmann-Scheffe. Având la dispoziție această teorie, este posibil, în multe situații, determinarea unui estimator MVU prin simpla studiere a densității de repartiție a caracteristicii studiate. Vom defini noțiunea de statistică suficientă ca fiind o statistică care sumarizează datele, în sensul că densitatea de repartiție a datelor nu mai depinde de parametrul necunoscut. Teorema de factorizare Neyman-Fisher ne va permite să găsim statistici suficiente prin examinarea densității de repartiție. Estimatorul MVU se poate determina prin găsirea unei funcții având ca argument statistica suficientă, funcție ce este un estimator nedeplasat. Pentru utilizarea apoi a Teoremei Rao-Blackwell-Lehmann-Scheffe trebuie să asigurăm faptul că statistica suficientă este completă, în sensul că există o unică funcție având-o ca argument ce este estimator nedeplasat.

2.2 Statistici suficiente

Problema estimării nivelului A a unui curent continuu situat într-un WGN revine la a arăta că media de selecție

$$\hat{A} = \bar{X} = \frac{1}{n} \sum_{i=0}^{n-1} X_i$$

este un estimator MVU , de dispersie σ^2/n . Dacă, pe de altă parte, am fi ales $\hat{A} = X_0$ drept estimatorul căutat, este evident faptul că, deși $\hat{A}$ este nedeplasat, dispersia sa este mult mai mare (este σ^2) decât minimul precedent. Intuitiv, performanța sa slabă este un rezultat al renunțării la variabilele de selecție $X_1, X_2, \dots, X_{n-1}$, ale căror realizări înmagazinează informații privind parametrul A . Se poate formula astfel întrebarea "care eșantioane sunt relevante pentru problema estimării". Următoarele mulțimi de date pot fi considerate a fi *suficiente* pentru determinarea estimatorului $\hat{A}$:

$$\begin{cases} S_1 &= \{x_0, x_1, \dots, x_{n-1}\} \\ S_2 &= \{x_0 + x_1, x_2, x_3, \dots, x_{n-1}\} \\ S_3 &= \left\{ \sum_{i=0}^{n-1} x_i \right\}, \end{cases}$$

unde, pentru fiecare i , x_i sunt datele observate, corespunzătoare variabilei de selecție X_i . Mulțimea S_1 reprezintă mulțimea originală de date, care, așa cum este de așteptat, este suficientă pentru problema estimării. Cum și S_2 și S_3 sunt și ele suficiente, arată că există mai multe mulțimi suficiente de date. Acea mulțime care conține numărul minim de elemente se numește *mulțimea minimală*. Dacă privim acum elementele acestor mulțimi drept statistici, spunem că cele n statistici ale mulțimii S_1 sunt suficiente, la fel și cele $n-1$ statistici ale mulțimii S_2 , la fel și unica statistică a mulțimii S_3 . Această unică statistică, $\sum_{i=0}^{n-1} X_i$, pe lângă faptul că este statistică suficientă, este chiar *statistica suficientă minimală*. Pentru estimarea parametrului A , odată ce știm $\sum_{i=0}^{n-1} X_i$, nu mai avem nevoie de datele (valorile) individuale ale fiecărei variabile de selecție, deoarece toate informațiile au fost încapsulate în statistica suficientă. Pentru a înțelege ce am spus prin afirmația anterioară, considerăm densitatea vectorului aleator de selecție (adică funcția de verosimilitate):

$$\mathcal{L}(v, A) = f(v, A) = \frac{1}{(2\pi\sigma^2)^{n/2}} \exp\left(-\frac{1}{2\sigma^2} \sum_{i=0}^{n-1} (x_i - A)^2\right) \quad (40)$$

și presupunem că au fost observate datele $T(v) = \sum_{i=0}^{n-1} x_i = T_0$. Cunoașterea valorii statisticii $\sum_{i=0}^{n-1} X_i$ va schimba densitatea de repartiție în cea condiționată:

$$f(v | \sum_{i=0}^{n-1} x_i = T_0, A),$$

care acum furnizează densitatea de repartiție a observațiilor după ce statistica suficientă a fost observată. Deoarece statistica este suficientă pentru estimarea parametrului A , această densitate condiționată ar trebui să nu depindă de parametrul estimat, A . Dacă ar depinde, ar trebui informații suplimentare asupra lui A .

Exemplul 2.1 Verificarea unei statistici suficiente. Considerăm densitatea vectorială de repartiție dată de (40). Pentru a arăta că $\sum_{i=0}^{n-1} X_i$ este o statistică suficientă, trebuie să determinăm $f(v | T(v) = T_0, A)$, unde $T(v) = \sum_{i=0}^{n-1} x_i$. Din definiția densității condiționate, avem:

$$f(v | T(v) = T_0, A) = \frac{f(v, T(v) = T_0, A)}{f(T(v) = T_0, A)} = \frac{f(v, A) \delta(T(v) - T_0)}{f(T(v) = T_0, A)}. \quad (41)$$

Pentru obținerea relației (41), am avut în vedere faptul că $T(v)$ este dependentă funcțional de v , deci densitatea comună $f(v, T(v) = T_0, A)$ este nenulă doar atunci când v satisface $T(v) = T_0$. Densitatea comună (numărătorul fracției) va fi deci $f(v, A) \delta(T(v) - T_0)$, δ fiind distribuția Dirac. Deoarece $T(v) \sim \mathcal{N}(nA, n\sigma^2)$, avem:

$$\begin{aligned} f(v, A) \delta(T(v) - T_0) &= \frac{1}{(2\pi\sigma^2)^{n/2}} \exp\left(-\frac{1}{2\sigma^2} \sum_{i=0}^{n-1} (x_i - A)^2\right) \delta(T(v) - T_0) \\ &= \frac{1}{(2\pi\sigma^2)^{n/2}} \exp\left(-\frac{1}{2\sigma^2} \left(\sum_{i=0}^{n-1} x_i^2 - 2AT(v) + nA^2\right)\right) \delta(T(v) - T_0) \\ &= \frac{1}{(2\pi\sigma^2)^{n/2}} \exp\left(-\frac{1}{2\sigma^2} \left(\sum_{i=0}^{n-1} x_i^2 - 2AT_0 + nA^2\right)\right) \delta(T(v) - T_0). \end{aligned}$$

Din (41),

$$\begin{aligned} f(v|T(v) = T_0, A) &= \frac{\frac{1}{(2\pi\sigma^2)^{n/2}} \exp\left(-\frac{1}{2\sigma^2} \sum_{i=0}^{n-1} x_i^2\right) \exp\left(-\frac{1}{2\sigma^2} (-2AT_0 + nA^2)\right)}{\frac{1}{\sqrt{2\pi n\sigma^2}} \exp\left(-\frac{1}{2n\sigma^2} (T_0 - nA)^2\right)} \delta(T(v) - T_0) \\ &= \frac{\sqrt{n}}{(2\pi\sigma^2)^{(n-1)/2}} \exp\left(-\frac{1}{2\sigma^2} \sum_{i=0}^{n-1} x_i^2\right) \exp\left(\frac{T_0^2}{2n\sigma^2}\right) \delta(T(v) - T_0), \end{aligned}$$

cantitate care nu depinde de A . Concluzionăm că $\sum_{i=0}^{n-1} X_i$ este o statistică suficientă pentru estimarea parametrului A .

Acest exemplu indică o metodă pentru a verifica că o statistică este suficientă. Însă, în multe situații, problema determinării densității condiționate este dificilă și, prin urmare, este nevoie de o abordare mai ușoară. De asemenea, în general, o altă chestiune, chiar mai importantă, care trebuie avută în vedere, este *identificarea potențialei statistici suficiente*. Un instrument important care ne permite rezolvarea acestei ultime chestiuni este reprezentat de *Teorema de factorizare Neyman-Fisher*.

Prezentăm, pentru început acest rezultat, după care, îl vom folosi pentru identificarea statisticilor suficiente în cadrul câtorva exemple reprezentative.

Teorema 2 (Teorema de factorizare Fisher-Neyman pentru parametru scalar) Dacă putem factoriza densitatea de repartiție

$$f(v, \theta) = g(T(v, \theta)) h(v), \quad (42)$$

unde g este o funcție dependentă de v doar prin intermediul lui $T(v)$, iar h este o funcție dependentă doar de v , atunci $T(V)$ este o statistică suficientă pentru parametrul θ . Reciproc, dacă $T(V)$ este o statistică suficientă pentru θ , atunci densitatea de repartiție a vectorului aleator de selecție (adică funcția de verosimilitate) poate fi factorizată precum în formula (42).

Demonstrație. Considerăm densitatea de repartiție comună (vectorială) $f(v, T(v), \theta)$. Pentru evaluarea acestei densități, menționăm că $T(v)$ este dependentă funcțional de v . Prin urmare, densitatea comună va fi zero atunci când este calculată în $v = v_0$, $T(v) = T_0$ cu excepția cazului când $T(v_0) = T_0$. Avem astfel

$$f(v, T(v) = T_0, \theta) = f(v, \theta) \delta(T(v) - T_0),$$

formulă ce va fi utilizată în cadrul demonstrației teoremei.

Un al doilea rezultat de care ne vom folosi constă în modul de a reprezenta densitatea de repartiție a unei funcții ce depinde de mai multe variabile aleatoare. Astfel, dacă $Y = g(X)$, unde X este un vector aleator n -dimensional, iar $g: \mathbb{R}^n \rightarrow \mathbb{R}$, atunci densitatea de repartiție a variabilei aleatoare Y poate fi scrisă astfel:

$$f(y) = \int \cdots \int_{\mathbb{R}^n} f(v) \delta(y - g(v)) dv \quad (43)$$

Această reprezentare este echivalentă cu formula clasică a transformării variabilelor aleatoare. Dacă, de exemplu X este o variabilă aleatoare 1-dimensională, atunci

$$\delta(y - g(x)) = \sum_{i=1}^k \frac{\delta(x - x_i)}{|g'(x)|_{x=x_i}}, \quad \text{unde } \{x_1, \dots, x_k\} \text{ sunt soluțiile ecuației } y = g(x).$$

Dacă înlocuim această formulă în (43) obținem formula clasică de transformare a variabilelor aleatoare.

Demonstrăm acum, pentru început, că $T(v)$ este o statistică suficientă, atunci când formula de factorizare are loc. Avem, datorită formulei de factorizare aplicată pentru deducerea ultimei egalități,

$$f(v|T(v) = T_0, \theta) = \frac{f(v, T(v) = T_0, \theta)}{f(T(v) = T_0, \theta)} = \frac{f(v, \theta) \delta(T(v) - T_0)}{f(T(v) = T_0, \theta)} = \frac{g(T(v) = T_0, \theta) h(v) \delta(T(v) - T_0)}{f(T(v) = T_0, \theta)} \quad (44)$$

Dar, în conformitate cu (43),

$$\begin{aligned} f(T(v) = T_0, \theta) &= \int \cdots \int_{\mathbb{R}^n} f(v, \theta) \delta(T(v) - T_0) dv = \int \cdots \int_{\mathbb{R}^n} g(T(v) = T_0, \theta) h(v) \delta(T(v) - T_0) dv \\ &= g(T(v) = T_0, \theta) \int \cdots \int_{\mathbb{R}^n} h(v) \delta(T(v) - T_0) dv. \end{aligned}$$

Ultima egalitate din formula anterioară are loc deoarece integrala este zero pe complementara suprafeței din $\mathbb{R}^n$ pentru care $T(v) = T_0$. Pe această suprafață, g este constantă. Inserăm formula obținută în (44) și rezultă:

$$f(v|T(v) = T_0, \theta) = \frac{h(v) \delta(T(v) - T_0)}{\int \cdots \int_{\mathbb{R}^n} h(v) \delta(T(v) - T_0) dv},$$

formulă ce nu depinde de parametrul θ . În concluzie, $T(V)$ este o statistică suficientă.

Pentru cea de a doua parte a teoremei, arătăm că, dacă $T(V)$ este o statistică suficientă, atunci formula de factorizare trebuie să aibă loc. Considerăm densitatea de repartiție comună

$$f(v, T(v) = T_0, \theta) = f(v|T(v) = T_0, \theta) f(T(v) = T_0, \theta) \quad (45)$$

și subliniem faptul că $f(v, T(v) = T_0, \theta) = f(v, \theta) \delta(T(v) - T_0)$. Deoarece $T(V)$ este o statistică suficientă, densitatea condiționată nu trebuie să depindă de parametrul θ și, prin urmare, o putem scrie ca $f(v|T(v) = T_0)$. Acum, deoarece avem $T(v) = T_0$ (ceea ce reprezintă o suprafață din $\mathbb{R}^n$), **densitatea condiționată este nevidă doar pe această suprafață**. Notăm $f(v|T(v) = T_0) = w(v) \delta(T(v) - T_0)$, unde

$$\int \cdots \int_{\mathbb{R}^n} w(v) \delta(T(v) - T_0) dv = 1. \quad (46)$$

Obținem, din (45),

$$f(v, \theta) \delta(T(v) - T_0) = w(v) \delta(T(v) - T_0) f(T(v) = T_0, \theta) \quad (47)$$

Definim acum

$$w(v) := \frac{h(v)}{\int \cdots \int_{\mathbb{R}^n} h(v) \delta(T(v) - T_0) dv},$$

astfel încât (46) să fie satisfăcută. Reprezentarea (47) devine:

$$f(v, \theta) \delta(T(v) - T_0) = \frac{h(v) \delta(T(v) - T_0)}{\int \cdots \int_{\mathbb{R}^n} h(v) \delta(T(v) - T_0) dv} f(T(v) = T_0, \theta),$$

sau, echivalent,

$$f(v, \theta) = g(T(v) = T_0, \theta) h(v), \quad \text{unde} \quad g(T(v) = T_0, \theta) := \frac{f(T(v) = T_0, \theta)}{\int \cdots \int_{\mathbb{R}^n} h(v) \delta(T(v) - T_0) dv}.$$

Scrișă sub această formă, densitatea de repartiție a statisticii suficiente poate fi determinată cu ajutorul formulei de factorizare, deoarece

$$f(T(v) = T_0, \theta) = g(T(v) = T_0, \theta) \int \cdots \int_{\mathbb{R}^n} h(v) \delta(T(v) - T_0) dv.$$

Demonstrația este, în acest moment, încheiată.

Trebuie menționat faptul că, uneori, nu este evident faptul că densitatea de repartiție poate fi factorizată sub forma dorită. În aceste situații, este posibil să nu existe o statistică suficientă. Prezentăm în continuare câteva exemple în care utilizăm instrumentul teoretic prezentat anterior.

Exemplul 2.2 Mărimea unui curent continuu aflat într-un WGN. Reexaminăm problema discutată în exemplul precedent. Densitatea era dată de formula (40), unde σ^2 era presupus a fi cunoscută. Pentru a demonstra că există o factorizare, observăm că:

$$\sum_{i=0}^{n-1} (x_i - A)^2 = \sum_{i=0}^{n-1} x_i^2 - 2A \sum_{i=0}^{n-1} x_i + nA^2,$$

ceea ce face ca densitatea să poată fi factorizată astfel:

$$f(v, A) = \underbrace{\frac{1}{(2\pi\sigma^2)^{n/2}} \exp\left(-\frac{1}{2\sigma^2} \left(nA^2 - 2A \sum_{i=0}^{n-1} x_i\right)\right)}_{=:g(T(v), A)} \underbrace{\exp\left(-\frac{1}{2\sigma^2} \sum_{i=0}^{n-1} x_i^2\right)}_{=:h(v)}.$$

Este evident că $T(V)$ este o statistică suficientă pentru parametrul A . Observăm că $T'(V) = 2 \sum_{i=0}^{n-1} X_i^2$ este, de asemenea, o statistică suficientă pentru A . Mai mult, se poate arăta că orice funcție bijectivă ce are ca argument pe $\sum_{i=0}^{n-1} X_i$ este o statistică suficientă pentru A . Prin urmare, statisticile suficiente sunt unice până la o transformare bijectivă.

Exemplul 2.3 Puterea unui WGN. Considerăm acum densitatea dată de (40), cu $A = 0$ și dispersia σ^2 fiind parametrul necunoscut. Avem astfel,

$$f(v, \sigma^2) = \frac{1}{(2\pi\sigma^2)^{n/2}} \exp\left(-\frac{1}{2\sigma^2} \sum_{i=0}^{n-1} x_i^2\right) \cdot 1 = g(T(v), A) h(v).$$

Din Teorema de factorizare, se poate arăta că $T(V) := \sum_{i=0}^{n-1} X_i^2$ este o statistică suficientă pentru parametrul σ^2 . Pentru aceasta, se arată că $f\left(v \mid \sum_{i=0}^{n-1} x_i^2 = T_0, \sigma^2\right)$ nu depinde de σ^2 , plecând de la faptul că densitatea de repartiție a variabilei aleatoare

$$S := \sum_{i=0}^{n-1} \frac{X_i^2}{\sigma^2} \sim \chi(n, 1), \quad \text{adică} \quad f_S(x) = \frac{1}{2^{n/2} \Gamma(n/2)} \exp(-s/2) s^{n/2-1} \mathbf{1}_{\{s \geq 0\}}.$$

Exemplul 2.4 Parametrul fazei unui curent sinusoidal. Reconsiderăm Exemplul 1.4, în care dorim să estimăm parametrul fazei unui curent de tip sinusoidal aflat într-un WGN:

$$X_i = A \cos(2\pi f_0 i + \phi) + W_i, \quad i \in \{0, 1, \dots, n-1\}.$$

Presupunem amplitudinea A și frecvența f_0 , precum și dispersia σ^2 a zgomotului alb cunoscute. Densitatea comună (funcția de verosimilitate) este:

$$\begin{aligned} f(v, \phi) &= \frac{1}{(2\pi\sigma^2)^{n/2}} \exp\left(-\frac{1}{2\sigma^2} \sum_{i=0}^{n-1} (x_i - A \cos(2\pi f_0 i + \phi))^2\right) \\ &= \frac{1}{(2\pi\sigma^2)^{n/2}} \exp\left(-\frac{1}{2\sigma^2} \left(\sum_{i=0}^{n-1} x_i^2 - 2A \sum_{i=0}^{n-1} x_i \cos(2\pi f_0 i + \phi) + \sum_{i=0}^{n-1} A^2 \cos^2(2\pi f_0 i + \phi)\right)\right) \\ &= \frac{1}{(2\pi\sigma^2)^{n/2}} \exp\left(-\frac{1}{2\sigma^2} \left(\sum_{i=0}^{n-1} x_i^2 - 2A \left(\sum_{i=0}^{n-1} x_i \cos(2\pi f_0 i)\right) \cos \phi + 2A \left(\sum_{i=0}^{n-1} x_i \sin(2\pi f_0 i)\right) \sin \phi \right.\right. \\ &\quad \left.\left. + \sum_{i=0}^{n-1} A^2 \cos^2(2\pi f_0 i + \phi)\right)\right). \end{aligned}$$

În acest exemplu nu apare faptul că densitatea este factorizabilă așa cum este nevoie în Teorema Neyman-Fisher, adică nu ar exista nicio statistică suficientă pentru parametrul fazei, ϕ . Cu toate acestea, putem scrie sub forma factorizată următoare:

$$f(v, \phi) = \underbrace{\frac{1}{(2\pi\sigma^2)^{n/2}} \exp\left(-\frac{1}{2\sigma^2} \left(\sum_{i=0}^{n-1} A^2 \cos^2(2\pi f_0 i + \phi) - 2AT_1(v) \cos \phi + 2AT_2(v) \sin \phi\right)\right)}_{=:g(T_1(v), T_2(v), \phi)} \underbrace{\exp\left(-\frac{1}{2\sigma^2} \sum_{i=0}^{n-1} x_i^2\right)}_{=:h(v)},$$

unde

$$T_1(v) := \sum_{i=0}^{n-1} x_i \cos(2\pi f_0 i) \quad \text{și} \quad T_2(v) := \sum_{i=0}^{n-1} x_i \sin(2\pi f_0 i).$$

O generalizare a Teoremei Neyman-Fisher afirmă că $(T_1(V), T_2(V))$ sunt, ca pereche, statistici suficiente pentru parametrul ϕ , fără a exista însă o statistică suficientă singulară. Motivul pentru care ne vom restrange atenția doar asupra statisticilor suficiente singulare va fi clarificat în secțiunea următoare.

Conceptul de statistici suficiente grupate trebuie privit astfel. Spunem că r statistici $T_1(V), T_2(V), \dots, T_r(V)$ formează o statistică suficientă grupată dacă densitatea condiționată $f(v, T_1(v), T_2(v), \dots, T_r(v), \theta)$ nu depinde de θ . Generalizarea Teoremei Neyman-Fisher spune că, dacă densitatea comună (a vectorului aleator de selecție) $f(v, \theta)$ poate fi factorizată sub forma (cititorul interesat poate consulta Kendall, Stuart [4]):

$$f(v, \theta) = g(T_1(v), T_2(v), \dots, T_r(v), \theta) h(v), \quad (48)$$

atunci r -uplul $\{T_1(V), T_2(V), \dots, T_r(V)\}$ formează o statistică suficientă grupată pentru parametrul θ .

Observația 2.1 Este evident faptul că datele observate furnizează întotdeauna o astfel de statistică grupată deoarece, dacă stabilim $r = n$ și $T_{i+1}(V) := X_i$, $i = 0, 1, \dots, n-1$, atunci putem alege $g = f(v, \theta)$ și $h = 1$, iar condiția (48) va fi identic satisfăcută. Desigur, foarte rar, ele formează mulțimea minimală de statistici suficiente.

2.3 Utilizarea suficienței în determinarea estimatorilor MVU

Presupunând că am reușit să determinăm o statistică suficientă $T(V)$ pentru parametrul θ , putem utiliza apoi *Teorema Rao-Blackwell-Lehmann-Scheffe (Teorema RBLs)* pentru a găsi un estimator MVU. Prezentăm, pentru început, un exemplu care să illustreze abordarea problemei, după care vom introduce cadrul teoretic de lucru.

Exemplul 2.5 Mărirea curentului continuu în WGN. Vom continua analiza Exemplului 2.2. Deși deja am obținut faptul că $\hat{A} = \bar{X}$ este un estimator MVU (deoarece este eficient), utilizăm acum *Teorema RBLs*, care poate fi folosită chiar și atunci când nu există un estimator eficient, deci metoda CRLB nu este viabilă. Procedura pentru determinarea estimatorului MVU, $\hat{A}$, poate fi aplicată în două moduri. Ambele se bazează pe statistica suficientă $T(V) = \sum_{i=0}^{n-1} X_i$:

1. Găsim orice estimator nedeplasat al parametrului A , să presupunem $\bar{A} = X_0$, și determinăm apoi $\hat{A} = \mathbb{E}(\bar{A}|T)$.
2. Determinăm o funcție g astfel încât $\hat{A} = g(T)$ este un estimator nedeplasat pentru parametrul A .

Pentru prima abordare putem pleca de la estimatorul nedeplasat $\bar{A} = X_0$ și determinăm $\hat{A} = \mathbb{E}(X_0|\sum_{i=0}^{n-1} X_i)$. Pentru a realiza aceasta avem nevoie de câteva proprietăți ale densității condiționate Gaussiene. Astfel, pentru un vector aleator Gaussian $(X, Y)^t$, cu media $\mu = (\mathbb{E}(X), \mathbb{E}(Y))^t$ și matricea de covarianță

$$C = \begin{pmatrix} D^2(X) & \text{cov}(X, Y) \\ \text{cov}(Y, X) & D^2(Y) \end{pmatrix},$$

densitatea de repartiție este

$$f_{(X,Y)}(x, y) = \frac{1}{2\pi \det^{1/2}(C)} \exp \left(-\frac{1}{2} \begin{pmatrix} x - \mathbb{E}(X) \\ y - \mathbb{E}(Y) \end{pmatrix}^t C^{-1} \begin{pmatrix} x - \mathbb{E}(X) \\ y - \mathbb{E}(Y) \end{pmatrix} \right).$$

Densitățile condiționate $f(x|y)$ și $f(y|x)$ sunt, de asemenea, Gaussiene și au loc următoarele formule pentru media și dispersia condiționată (demonstrațiile pot fi studiate în Kay [2, Theorem 10.1 și Theorem 10.2]):

$$\begin{cases} \mathbb{E}(X|Y) &= \int_{-\infty}^{+\infty} x f(x|y) dx = \int_{-\infty}^{+\infty} x \frac{f_{(X,Y)}(x, y)}{f_Y(y)} dx = \mathbb{E}(X) + \frac{\text{cov}(X, Y)}{D^2(Y)} (Y - \mathbb{E}(Y)). \\ D^2(X|Y) &= D^2(X) - \frac{\text{cov}^2(X, Y)}{D^2(Y)}. \end{cases} \quad (49)$$

Aplicăm acum aceste formule pentru variabilele aleatoare $X = X_0$ și $Y = \sum_{i=0}^{n-1} X_i$. Observăm, de asemenea, că:

$$\begin{pmatrix} X \\ Y \end{pmatrix} = \begin{pmatrix} X_0 \\ \sum_{i=0}^{n-1} X_i \end{pmatrix} = \underbrace{\begin{pmatrix} 1 & 0 & 0 & \dots & 0 \\ 1 & 1 & 1 & \dots & 1 \end{pmatrix}}_{=:L} \begin{pmatrix} X_0 \\ X_1 \\ \vdots \\ X_{n-1} \end{pmatrix}.$$

Prin urmare,

$$(X, Y) \sim \mathcal{N}(\mu, C), \quad \text{unde}$$

$$\mu = L\mathbb{E}(X) = LA\mathbf{1} = \begin{pmatrix} A \\ nA \end{pmatrix} \quad \text{și} \quad C = \sigma^2 LL^t = \sigma^2 \begin{pmatrix} 1 & 1 \\ 1 & n \end{pmatrix},$$

deoarece reprezintă o transformare liniară a unui vector aleator Gaussian. Formula (49) conduce la

$$\hat{A} = \mathbb{E}(X|Y) = A + \frac{\sigma^2}{n\sigma^2} \left(\sum_{i=0}^{n-1} X_i - nA \right) = \frac{1}{n} \sum_{i=0}^{n-1} X_i,$$

care este un estimator MVU. Această abordare, care presupune evaluarea mediei condiționate, este, de regulă, greu abordabilă din punct de vedere matematic.

În ceea ce privește cea de a doua abordare, trebuie să găsim o funcție g astfel încât

$$\hat{A} = g \left(\sum_{i=0}^{n-1} X_i \right)$$

să fie un estimator nedeplasat pentru A . Dacă considerăm $g(x) = x/n$, obținem imediat un estimator MVU. Această metodă este mai ușor de aplicat și este mai des utilizată în practică.

Teorema 3 (Rao-Blackwell-Lehmann-Scheffe, cazul scalar) Dacă $\bar{\theta}$ este un estimator nedeplasat pentru parametrul θ și $T(V)$ este o statistică suficientă pentru θ , atunci $\hat{\theta} = \mathbb{E}(\bar{\theta}|T(V))$ este:

1. un estimator valid pentru θ (adică nu depinde de θ);
2. nedeplasat;
3. de dispersie mai mică sau egală cu dispersia lui $\bar{\theta}$, pentru toți θ .

În plus, dacă statistica suficientă este completă, atunci $\hat{\theta}$ este estimatorul MVU.

Demonstrație. Vom analiza cazul parametrului 1-dimensional, θ . Pentru a demonstra (1), observăm că:

$$\hat{\theta} = \mathbb{E}(\bar{\theta}|T(V)) = \int \cdots \int_{\mathbb{R}^n} \bar{\theta}(v) f(v|T(v), \theta) dv \quad (50)$$

Din definiția unei statistici suficiente, $f(v|T(v), \theta)$ nu depinde de θ , ci de v și de $T(v)$. Prin urmare, în urma integrării, rezultatul va fi o funcție de T și nu va depinde de θ .

Pentru punctul (2), adică $\hat{\theta}$ este nedeplasat, din (50) și din faptul că $\bar{\theta}$ este dependent doar de T , avem:

$$\begin{aligned} \mathbb{E}(\hat{\theta}) &= \int \cdots \int_{\mathbb{R}^{n+1}} \bar{\theta}(v) f(v|T(v), \theta) dv f(T(v), \theta) dT = \int \cdots \int_{\mathbb{R}^n} \bar{\theta}(v) \int_{-\infty}^{+\infty} f(v|T(v), \theta) f(T(v), \theta) dT dv \\ &= \int \cdots \int_{\mathbb{R}^n} \bar{\theta}(v) f(v, \theta) dv = \mathbb{E}(\bar{\theta}) = \theta \quad \text{deoarece } \bar{\theta} \text{ este estimator nedeplasat.} \end{aligned}$$

Pentru a demonstra punctul (3), și anume că $D^2(\hat{\theta}) \leq D^2(\bar{\theta})$, procedăm astfel:

$$D^2(\bar{\theta}) = \mathbb{E}((\bar{\theta} - \mathbb{E}(\bar{\theta}))^2) = \mathbb{E}((\bar{\theta} - \hat{\theta} + \hat{\theta} - \theta)^2) = \mathbb{E}((\bar{\theta} - \hat{\theta})^2) + 2\mathbb{E}((\bar{\theta} - \hat{\theta})(\hat{\theta} - \theta)) + D^2(\hat{\theta}).$$

Știm că $\hat{\theta}$ este o funcție ce depinde doar de T . De aici rezultă că:

$$\mathbb{E}_{T, \bar{\theta}}((\bar{\theta} - \hat{\theta})(\hat{\theta} - \theta)) = \mathbb{E}_T \mathbb{E}_{\bar{\theta}|T}((\bar{\theta} - \hat{\theta})(\hat{\theta} - \theta)) = \mathbb{E}_T \mathbb{E}_{\bar{\theta}|T}(\bar{\theta} - \hat{\theta})(\hat{\theta} - \theta) = \mathbb{E}_T \underbrace{(\mathbb{E}_{\bar{\theta}|T}(\bar{\theta}|T) - \hat{\theta})}_{=\hat{\theta}-\hat{\theta}=0}(\hat{\theta} - \theta) = 0.$$

Am obținut astfel rezultatul dorit, adică

$$D^2(\bar{\theta}) = \mathbb{E}((\bar{\theta} - \hat{\theta})^2) + D^2(\hat{\theta}) \geq D^2(\hat{\theta}).$$

O statistică este completă dacă există o singură funcție ce depinde de statistică și care este nedeplasată. Justificăm acum că $\hat{\theta} = \mathbb{E}(\bar{\theta}|T(V))$ este estimatorul MVU. Considerăm toți estimatorii nedeplasați ai lui θ . Determinând $\mathbb{E}(\bar{\theta}|T(V))$ putem micșora dispersia estimatorului (punctul (3)) și rămânem în clasa estimatorilor admisibili (punctul (2)). Dar $\mathbb{E}(\bar{\theta}|T(V))$ este o funcție ce depinde doar de statistica suficientă $T(V)$ deoarece

$$\hat{\theta} = \mathbb{E}(\bar{\theta}|T(V)) = \int \cdots \int_{\mathbb{R}^n} \bar{\theta}(v) f(\bar{\theta}|T(v)) d\bar{\theta} = g(T(v)) \quad (51)$$

Dacă $T(V)$ este complet, există doar o singură funcție de T ce este estimator nedeplasat. Prin urmare, $\hat{\theta}$ este unic, independent de $\bar{\theta}$ ales din clasa estimatorilor nedeplasați. Fiecare $\bar{\theta}$ va furniza același estimator $\hat{\theta}$. Cum dispersia lui $\hat{\theta}$ este mai mică decât dispersia oricărui estimator $\bar{\theta}$ din clasă (punctul (3)), rezultă că $\hat{\theta}$ este estimatorul MVU. Sumarizând, estimatorul MVU se poate obține considerând orice estimator nedeplasat și aplicându-i transformarea dată de (51). Alternativ, cum există o unică funcție dependentă de statistica suficientă ce duce la un estimator nedeplasat, este suficient să o identificăm. Demonstrația este încheiată.

Pentru ultima afirmație din Teorema 3, am văzut în Exemplul 2.5 că $g(\sum_{i=0}^{n-1} X_i) = \sum_{i=0}^{n-1} X_i/n$. Proprietatea de completitudine depinde de densitatea de repartiție a vectorului de selecție V , ce determină densitatea statisticii suficiente. În general, este dificil de arătat că o statistică suficientă este completă, un studiu detaliat al acestei problematice putând fi analizat în Kendall, Stuart [4]. Vom prezenta în continuare două exemple, pentru a familiariza cititorul cu conceptul de completitudine al statisticilor suficiente.

Exemplul 2.6 Completitudinea statisticii suficiente. Pentru estimarea parametrului A , statistica suficientă $T = \sum_{i=0}^{n-1} X_i$ este completă, în sensul că există o unică funcție g pentru care $\mathbb{E}(g(T)) = A$. Presupunem, prin reducere la absurd, că ar mai exista încă o funcție h , cu aceeași proprietate: $\mathbb{E}(h(T)) = A$. Rezultă deci că

$$\mathbb{E}(g(T) - h(T)) = A - A = 0, \quad \text{pentru toți } A.$$

Deoarece $T \sim \mathcal{N}(nA, n\sigma^2)$, egalitatea de mai sus înseamnă:

$$\int_{-\infty}^{+\infty} u(T) \frac{1}{\sqrt{2\pi n\sigma^2}} \exp\left(-\frac{1}{2n\sigma^2} (T - nA)^2\right) dT = 0, \quad \text{pentru toți } A,$$

unde am notat $u(T) = g(T) - h(T)$. Dacă notăm $\tau = T/n$ și $w(\tau) = u(n\tau)$, avem:

$$\int_{-\infty}^{+\infty} w(\tau) \frac{n}{\sqrt{2\pi n\sigma^2}} \exp\left(-\frac{n}{2\sigma^2} (A - \tau)^2\right) d\tau = 0, \quad \text{pentru toți } A, \quad (52)$$

formulă care reprezintă convoluția funcției $w(\tau)$ cu un puls Gaussian, $w(\tau)$. Cum egalitatea (52) are loc pentru toți A , $w(\tau)$ trebuie să fie identic zero. De aici rezultă că $0 = u(n\tau) = g(n\tau) - h(n\tau)$, pentru orice τ , adică $g \equiv h$.

Exemplul 2.7 Lipsa de completitudine a statisticii suficiente. Considerăm estimatorul lui A dat de:

$$X_0 = A + W_0, \quad \text{unde } W_0 \sim \mathcal{U}\left(-\frac{1}{2}, \frac{1}{2}\right).$$

O statistică suficientă este X_0 , furnizând unica dată disponibilă și, prin urmare, X_0 este un estimator nedeplasat pentru A . Am putea concluziona că $g(X_0) = X_0$ este un candidat viabil pentru estimatorul MVU. Ar trebui verificat că este o statistică suficientă completă. Similar exemplului anterior, presupunem că există încă o funcție h ce verifică $\mathbb{E}(h(X_0)) = A$ și încercăm să arătăm că $h \equiv g$. Notăm iar $u(T) = g(T) - h(T)$ și căutăm posibilele soluții v pentru ecuația:

$$\int_{-\infty}^{+\infty} u(T) f(v, A) dv = 0, \quad \text{pentru toți } A.$$

Pentru această problemă, $V = X_0 = T$ și rezultă

$$\int_{-\infty}^{+\infty} u(T) f(T, A) dT = 0 \iff \int_{A-1/2}^{A+1/2} u(T) dT = 0, \quad \text{pentru toți } A,$$

deoarece

$$f(T, A) = \begin{cases} 1, & A - 1/2 \leq T \leq A + 1/2 \\ 0, & \text{în caz contrar.} \end{cases}$$

O soluție a ecuației de mai sus este $u(T) = \sin(2\pi T)$, iar de aici rezultă că $\sin(2\pi T) = g(T) - h(T)$, sau, echivalent, $h(T) = g(T) - \sin(2\pi T)$. Obținem astfel că estimatorul

$$\hat{A} = X_0 - 2 \sin(2\pi X_0)$$

este, de asemenea, bazat pe statistica suficientă și este un estimator nedeplasat pentru parametrul A . Cum am găsit cel puțin încă un estimator nedeplasat ce e funcție de statistica suficientă, putem concluziona că statistica suficientă nu este completă. Teorema RBLS nu mai poate fi aplicată și nu putem trage concluzia că $\hat{A} = X_0$ este estimatorul MVU.

Pentru a concluziona, prezentă etapele ce trebuie parcurse pentru determinarea unui estimator MVU:

1. Determinăm o statistică suficientă pentru parametrul θ , notată cu $T(V)$, folosind Teorema de factorizare Neman-Fisher.
2. Determinăm dacă statistica suficientă este completă, iar dacă da, putem continua. În caz contrar, această abordare nu poate fi utilizată.
3. Determinăm o funcție g , ce are ca argument statistica suficientă și care produce un estimator nedeplasat, $\hat{\theta} = g(T(V))$. Estimatorul MVU va fi $\hat{\theta}$. O variantă alternativă a acestui pas o constituie calcularea $\hat{\theta} = \mathbb{E}(\hat{\theta}|T(V))$, unde $\hat{\theta}$ este un estimator nedeplasat oarecare. Această ultimă abordare este, de regulă, greu de aplicat în practică, tocmai datorită modului de evaluare a mediei condiționate.

Exemplul 2.8 Valoarea medie a zgomotului uniform. Presupunem că avem un model în care datele observate x_i sunt corespunzătoare variabilelor de selecție $X_i = W_i$, unde $W_i, i \in \{0, 1, \dots, n-1\}$, sunt zgomote independente stochastice, identic repartizate $\mathcal{U}(0, \beta)$, $\beta > 0$. Ne propunem să determinăm un estimator MVU pentru media $\theta = \beta/2$.

Abordarea prin intermediul CRLB pentru determinarea unui estimator eficient și, prin urmare, care să fie și un estimator MVU nu poate fi utilizată pentru această situație deoarece densitatea de repartiție a vectorului aleator de selecție, adică funcția de verosimilitate, nu satisface condițiile de regularitate, având:

$$\mathbb{E} \left(\frac{\partial \ln \mathcal{L}(v, \theta)}{\partial \theta} \right) \neq 0, \quad \text{pentru toți } \theta.$$

Un estimator natural pentru θ este media de selecție $\bar{X} = (\sum_{i=0}^{n-1} X_i)/n$. Acesta este nedeplasat și are dispersia $D^2(\bar{X}) = \beta^2/(12n)$. Pentru a determina dacă acest estimator este MVU vom urma pașii descriși înaintea acestui exemplu. Construim, pentru început, funcția în scară

$$u(x) = \mathbf{1}_{\{x \geq 0\}} = \begin{cases} 1, & \text{dacă } x \geq 0, \\ 0, & \text{dacă } x < 0. \end{cases}$$

Obținem astfel densitatea fiecărei variabile aleatoare de selecție, precum și pentru funcția de verosimilitate:

$$\begin{cases} f(x_i, \theta) = \frac{1}{\beta} (u(x_i) - u(x_i - \beta)), & \text{unde } \beta = 2\theta, \\ \mathcal{L}(v, \theta) = \frac{1}{\beta^n} \prod_{i=0}^{n-1} (u(x_i) - u(x_i - \beta)) \neq 0 & \text{doar dacă } 0 < x_i < \beta, \text{ pentru toți indicii } i. \end{cases}$$

Astfel, obținem următoarea factorizare a densității vectoriale:

$$\mathcal{L}(v, \theta) = \mathcal{L}(x_0, \dots, x_{n-1}, \theta) = \mathbf{1}_{\{0 < x_i < \beta, \forall i\}} = \underbrace{\beta^{-n} u(\beta - \max_{i=0, n-1} x_i)}_{=: g(T(v), \theta)} \underbrace{u(\min_{i=0, n-1} x_i)}_{=: h(v)}.$$

Apelăm la Teorema de factorizare Neyman-Fisher și obținem că ultima statistică de ordine $T(V) = \max_{i=0, n-1} X_i$ este o statistică suficientă pentru parametrul θ .

În următoarea etapă determinăm o funcție ce are ca argument statistica suficientă și care produce o statistică nedeplasată. Calculăm, pentru început, repartiția statisticii de ordine T :

$$F_T(\xi) = \mathbb{P}(T \leq \xi) = \mathbb{P}(X_0 \leq \xi, \dots, X_{n-1} \leq \xi) = \prod_{i=0}^{n-1} \mathbb{P}(X_i \leq \xi) = (\mathbb{P}(X_0 \leq \xi))^n.$$

Densitatea de repartiție a lui T se obține prin derivarea funcției de repartiție calculate anterior:

$$\begin{aligned} f_T(\xi) &= n (\mathbb{P}(X_0 \leq \xi))^{n-1} \frac{d\mathbb{P}(X_0 \leq \xi)}{d\xi} = n (\mathbb{P}(X_0 \leq \xi))^{n-1} f_{X_0}(\xi) = n (\mathbb{P}(X_0 \leq \xi))^{n-1} \frac{1}{\beta} \mathbf{1}_{\{0 < \xi < \beta\}}. \\ &= n \left(\frac{\xi}{\beta} \right)^{n-1} \frac{1}{\beta} \mathbf{1}_{\{0 < \xi < \beta\}}. \end{aligned}$$

Media statisticii suficiente T este:

$$\mathbb{E}(T) = \int_{-\infty}^{+\infty} \xi f_T(\xi) d\xi = \int_0^\beta \xi n \left(\frac{\xi}{\beta} \right)^{n-1} \frac{1}{\beta} d\xi = \frac{n}{n+1} \beta = \frac{2n}{n+1} \theta.$$

Pentru a face acest estimator nedeplasat, trebuie ca $2n/(n+1) = 1$ și vom considera estimatorul $\hat{\theta} := ((n+1)/(2n)) T$, ceea ce spune că estimatorul MVU este

$$\hat{\theta} = \frac{n+1}{2n} \max_{i=0, n-1} X_i.$$

Cumva, contrar intuiției, **media de selecție nu este estimatorul MVU** pentru media zgomotului distribuit uniform. Dacă dorim să comparăm dispersia minimală cu cea a mediei de selecție, observăm că:

$$D^2(\hat{\theta}) = \left(\frac{n+1}{2n} \right)^2 D^2(T) = \left(\frac{n+1}{2n} \right)^2 \left(\int_0^\beta \xi^2 \frac{n\xi^{n-1}}{\beta^n} d\xi - \left(\frac{n\beta}{n+1} \right)^2 \right) = \left(\frac{n+1}{2n} \right)^2 \frac{n\beta^2}{(n+1)^2(n+2)} = \frac{\beta^2}{4n(n+2)}.$$

Se observă că, pentru $n \geq 2$,

$$D^2(\bar{X}) = \frac{\beta^2}{12n} > \frac{\beta^2}{4n(n+2)} = D^2(\hat{\theta}).$$

Succesul acestui exemplu se bazează pe posibilitatea de a determina o statistică suficientă singulară. Dacă ne referim la Exemplul 2.4, a estimării parametrului fazei, acolo am avut nevoie de două statistici $T_1(V)$ și $T_2(V)$, fapt ce face abordarea prezentată dificil de aplicat.

2.4 Extinderea la cazul parametrilor vectoriali

Extindem studiile precedente și intenționăm să determinăm un estimator vectorial MVU , de tipul $p \times 1$, pentru un parametru vectorial de aceeași dimensiune. Fiecare componentă a estimatorului trebuie să fie nedeplasată raportată la parametrul estimat și să aibă dispersie (varianță) minimă. În cazul vectorial, am putea avea mai multe statistici suficiente decât numărul de parametri estimați ($r > p$), exact tot atâtea ($r = p$), sau chiar mai puține ($r < p$). În exemplele abordate în această secțiune ne vom referi doar la primele două situații. O situație de dorit ar fi ca $r = p$, deoarece aceasta ar permite determinarea estimatorilor MVU prin transformarea statisticilor suficiente, păstrând proprietatea de nedeplasare.

O statistică vectorială $\mathbf{T}(V) = (T_1(V), T_2(V), \dots, T_r(V))^t$ se numește *suficientă* pentru estimarea parametrului vectorial θ dacă densitatea condiționată a vectorului de selecție în raport cu statistica, adică $f(V|\mathbf{T}(V), \theta)$, nu depinde de parametrul θ . Mai precis, aceasta înseamnă că $f(V|\mathbf{T}(V), \theta) = f(V|\mathbf{T}(V))$. Cum, în general, pot exista mai multe astfel de statistici $\mathbf{T}(V)$, suntem interesați de *statistica suficientă minimală*, sau, altfel spus, de *statistica $\mathbf{T}$ de dimensiune minimală*. Pentru aceasta, facem apel tot la Teorema de factorizare Neyman-Fisher, în varianta sa vectorială.

Teorema 4 (Teorema de factorizare Fisher-Neyman pentru parametru vectorial) Dacă putem factoriza densitatea de repartiție a vectorului aleator de selecție sub forma:

$$\mathcal{L}(v, \theta) = g(\mathbf{T}(v), \theta) h(v), \quad (53)$$

unde g este o funcție cu valori în $\mathcal{M}_{r \times 1}$, ce depinde doar de θ și de v prin intermediul lui $\mathbf{T}(v)$ și h este o funcție ce depinde doar de v , atunci $\mathbf{T}(V)$ este o statistică suficientă pentru θ . Reciproc, dacă $\mathbf{T}(V)$ este o statistică suficientă pentru θ , funcția de verosimilitate poate fi factorizată precum în formula (53).

Demonstrația, pe care o vom omite aici, poate fi studiată în Kendall, Stuart [4]. Precizăm că, întotdeauna, va exista o mulțime de statistici suficiente, mulțimea de date oferite de V fiind, chiar ele, suficiente. Problema care trebuie cercetată este dimensiunea statisticii suficiente, precum și concluzia oferită de Teorema 4.

Exemplul 2.9 Estimarea ansamblului parametrilor unui curent sinusoidal. Presupunem că avem un curent de tip sinusoidal, aflat într-un WGN, pentru care datele observate sunt date de variabilele aleatoare de selecție:

$$X_i = A \cos(2\pi f_0 i) + W_i, \quad i \in \{0, 1, \dots, n-1\},$$

unde amplitudinea A , frecvența f_0 și dispersia zgomotului, σ^2 , sunt parametri necunoscuți. Vectorul parametrilor va fi $\theta = (A, f_0, \sigma^2)^t$, iar densitatea comună a datelor este:

$$\begin{aligned} \mathcal{L}(v, \theta) &= \frac{1}{(2\pi\sigma^2)^{n/2}} \exp \left(-\frac{1}{2\sigma^2} \sum_{i=0}^{n-1} (x_i - A \cos(2\pi f_0 i))^2 \right) \\ &= \frac{1}{(2\pi\sigma^2)^{n/2}} \exp \left(-\frac{1}{2\sigma^2} \left(\sum_{i=0}^{n-1} x_i^2 - 2A \sum_{i=0}^{n-1} x_i \cos(2\pi f_0 i) + \sum_{i=0}^{n-1} A^2 \cos^2(2\pi f_0 i) \right) \right). \end{aligned}$$

Datorită termenului $\sum_{i=0}^{n-1} x_i \cos(2\pi f_0 i)$, cu f_0 necunoscut, nu putem reduce densitatea anterioară la forma dată de (53). Dacă, în schimb, frecvența ar fi cunoscută, vectorul parametrilor necunoscuți este $\theta = (A, \sigma^2)^t$, iar densitatea se poate reprezenta imediat astfel:

$$\mathcal{L}(v, \theta) = \frac{1}{(2\pi\sigma^2)^{n/2}} \exp \left(-\frac{1}{2\sigma^2} \underbrace{\left(\sum_{i=0}^{n-1} x_i^2 - 2A \sum_{i=0}^{n-1} x_i \cos(2\pi f_0 i) + \sum_{i=0}^{n-1} A^2 \cos^2(2\pi f_0 i) \right)}_{=:g(\mathbf{T}(v), \theta)} \right) \underbrace{1}_{=:h(v)},$$

unde

$$\mathbf{T}(v) := \begin{pmatrix} \sum_{i=0}^{n-1} x_i \cos(2\pi f_0 i) & \sum_{i=0}^{n-1} x_i^2 \end{pmatrix}^t.$$

Statistica $T(V)$ este o statistică suficientă pentru A și σ^2 , dar doar dacă f_0 este cunoscut.

Presupunem acum că avem:

$$X_i = A + W_i, \quad \text{unde } W_i \text{ este WGN, pentru } i \in \{0, 1, \dots, n-1\}$$

și presupunem că amplitudinea A și dispersia σ^2 a zgomotului alb Gaussian sunt parametri necunoscuți. Putem folosi rezultatele obținute în prima parte a exemplului pentru a determina o statistică suficientă de dimensiune 2 pentru estimarea mărimii curentului continuu și a dispersiei zgomotului. Parametrul necunoscut este $\theta = (A, \sigma^2)^t$, iar statistica suficientă este:

$$T(V) = \begin{pmatrix} \sum_{i=0}^{n-1} X_i & \sum_{i=0}^{n-1} X_i^2 \end{pmatrix}^t. \quad (54)$$

Am văzut că, dacă σ^2 este cunoscut, $\sum_{i=0}^{n-1} X_i$ este o statistică suficientă pentru amplitudinea A (vezi Exemplul 2.2), iar dacă A este cunoscut ($A = 0$), $\sum_{i=0}^{n-1} X_i^2$ este o statistică suficientă pentru σ^2 (vezi Exemplul 2.3). În analiza de față, aceleași statistici sunt, în ansamblu, ca pereche, suficiente. Acest lucru nu este, întotdeauna, adevărat.

Cum, în exemplul precedent, avem doi parametri și am obținut două statistici suficiente, putem acum determina un estimator MVU pentru parametrul vectorial. Pentru a realiza aceasta, enunțăm varianta vectorială a Teoremei Rao-Blackwell-Lehmann-Scheffe.

Teorema 5 (Rao-Blackwell-Lehmann-Scheffe, cazul vectorial) Dacă $\bar{\theta}$ este un estimator nedeplasat pentru parametrul θ și $\mathbf{T}(V)$ este o statistică suficientă, $r \times 1$ dimensională, pentru θ , atunci $\hat{\theta} = \mathbb{E}(\bar{\theta} | \mathbf{T}(V))$ este:

1. un estimator valid pentru θ (adică nu depinde de θ);
2. nedeplasat;
3. de dispersie mai mică sau egală cu dispersia lui $\bar{\theta}$ (compararea făcându-se componentă cu componentă), pentru toți θ .

În plus, dacă statistica suficientă este completă, atunci $\hat{\theta}$ este estimatorul MVU.

În cazul vectorial, completitudinea înseamnă că pentru o funcție, $r \times 1$ -dimensională, de $\mathbf{T}$, $\mathbf{u}(\mathbf{T})$, dacă

$$\mathbb{E}(\mathbf{u}(\mathbf{T})) = \int \cdots \int_{\mathbb{R}^r} \mathbf{u}(\mathbf{T}) \mathcal{L}(\mathbf{T}, \theta) d\mathbf{T} = \mathbf{0}, \quad \text{pentru toți } \theta, \quad (55)$$

atunci $\mathbf{u}(\mathbf{T}) = \mathbf{0} \in \mathcal{M}_{r \times 1}(\mathbb{R})$, pentru toți $\mathbf{T}$. Similar cazului 1-dimensional pentru parametru, condiția este greu verificabil în practică. Pentru a determina estimatorul MVU direct cu $\mathbf{T}(V)$, fără a mai evalua $\mathbb{E}(\bar{\theta} | \mathbf{T}(V))$, dimensiunea statisticii suficiente ar trebui să fie egală cu dimensiunea vectorului parametrilor necunoscuți, $r = p$. Dacă această condiție este verificată, ar trebui doar să găsim o funcție p -dimensională $\mathbf{g}$, astfel încât

$$\mathbb{E}(\mathbf{g}(\mathbf{T})) = \theta,$$

în ipoteza că $\mathbf{T}$ este o statistică suficientă completă.

Exemplul 2.10 (continuarea Exemplului 2.9) Pentru statistica suficientă (54), media sa este:

$$\mathbb{E}(\mathbf{T}(V)) = \begin{pmatrix} \mathbb{E}(T_1(V)) & \mathbb{E}(T_2(V)) \end{pmatrix} = \begin{pmatrix} nA & n\mathbb{E}(X_0^2) \end{pmatrix}^t = \begin{pmatrix} nA & n(\sigma^2 + A^2) \end{pmatrix}^t.$$

Desigur, am fi putut rezolva problema deplasării pentru prima componentă a lui $\mathbf{T}(V)$, prin împărțirea sa la n , dar ar fi dăunat celei de a doua componente. $T_2(V)$ estimează momentul teoretic de ordinul doi și nu dispersia dorită. Dacă transformăm pe $\mathbf{T}(V)$ astfel:

$$\mathbf{g}(\mathbf{T}(V)) = \begin{pmatrix} \frac{1}{n} T_1(V) \\ \frac{1}{n} T_2(V) - \left(\frac{1}{n} T_1(V) \right)^2 \end{pmatrix} = \begin{pmatrix} \bar{X} \\ \frac{1}{n} \sum_{i=0}^{n-1} X_i^2 - \bar{X}^2 \end{pmatrix},$$

atunci, deoarece $\bar{X} \sim \mathcal{N}(A, \sigma^2/n)$,

$$\mathbb{E}(\bar{X}) = A \quad \text{și} \quad \mathbb{E}\left(\frac{1}{n} \sum_{i=0}^{n-1} X_i^2 - \bar{X}^2\right) = \sigma^2 + A^2 - \mathbb{E}(\bar{X}^2) = \sigma^2 + A^2 - A^2 - \frac{\sigma^2}{n} = \frac{n-1}{n} \sigma^2.$$

Multiplîcând atunci statistica cu $n/(n-1)$, va fi nedeplasată pentru parametrul σ^2 . Transformarea recomandată va fi:

$$\mathbf{g}(\mathbf{T}(V)) = \begin{pmatrix} \frac{1}{n} T_1(V) \\ \frac{1}{n-1} \left(T_2(V) - n \left(\frac{1}{n} T_1(V) \right)^2 \right) \end{pmatrix} = \begin{pmatrix} \bar{X} \\ \frac{1}{n-1} \left(\sum_{i=0}^{n-1} X_i^2 - n \bar{X}^2 \right) \end{pmatrix}.$$

Un simplu calcul algebric arată că:

$$\sum_{i=0}^{n-1} (X_i - \bar{X})^2 = \sum_{i=0}^{n-1} X_i^2 - 2 \sum_{i=0}^{n-1} X_i \bar{X} + n \bar{X}^2 = \sum_{i=0}^{n-1} X_i^2 - n \bar{X}^2,$$

iar estimatorul va fi:

$$\hat{\theta} = \begin{pmatrix} \bar{X} \\ \frac{1}{n-1} \sum_{i=0}^{n-1} (X_i - \bar{X})^2 \end{pmatrix},$$

care este un estimator MVU pentru $\theta = (A, \sigma^2)^t$. Factorul de normalizare $1/(n-1)$ pentru $\hat{\sigma}^2$ se datorează pierderii unui grad de libertate în estimarea mediei. În acest exemplu, $\hat{\theta}$ nu este un estimator eficient. Se poate demonstra (vom omite aici această demonstrație) că $\bar{X}$ și $\hat{\sigma}^2$ sunt statistici independente și că

$$\bar{X} \sim \mathcal{N}\left(A, \frac{\sigma^2}{n}\right) \quad \text{și} \quad \frac{(n-1)\hat{\sigma}^2}{\sigma^2} \sim \chi^2(n-1).$$

Matricea de covarianță și CRLB (vezi Exemplul 1.5) sunt:

$$C_{\hat{\theta}} = \begin{pmatrix} \frac{\sigma^2}{n} & 0 \\ 0 & \frac{2\sigma^4}{n-1} \end{pmatrix} \quad \text{și} \quad I^{-1}(\theta) = \begin{pmatrix} \frac{\sigma^2}{n} & 0 \\ 0 & \frac{2\sigma^4}{n} \end{pmatrix}.$$

Se observă că estimatorul MVU nu se putea obține prin examinarea CRLB.

În final, observăm că dacă am fi știut sau am fi intuit că media de selecție și că dispersia de selecție modificată sunt estimatori MVU, am fi putut simplifica studiul. Densitatea vectorului de selecție este:

$$\begin{aligned} \mathcal{L}(v, \theta) &= \frac{1}{(2\pi\sigma^2)^{n/2}} \exp\left(-\frac{1}{2\sigma^2} \sum_{i=0}^{n-1} (x_i - A)^2\right) \\ &= \underbrace{\frac{1}{(2\pi\sigma^2)^{n/2}} \exp\left(-\frac{1}{2\sigma^2} \left(\sum_{i=0}^{n-1} (x_i - \bar{x})^2 + n(\bar{x} - A)^2\right)\right)}_{=: \mathbf{g}(\mathbf{T}'(V), \theta)} \cdot \underbrace{1}_{=: h(v)}, \end{aligned}$$

unde $\mathbf{T}'(V) = (\bar{X}, \sum_{i=0}^{n-1} (X_i - \bar{X})^2)^t$. Împărțind a doua componentă la $(n-1)$ obținem estimatorul $\hat{\theta}$. Desigur, există o transformare bijectivă între $\mathbf{T}'(V)$ și $\mathbf{T}(V)$, ceea ce subliniază, încă o dată, că statisticile suficiente sunt unice până la o transformare bijectivă.

Pentru concluzia că $\hat{\theta}$ este un estimator MVU nu am verificat completitudinea statisticii. Aceasta rezultă datorită faptului că densitatea Gaussiană este un caz particular al familiilor de densități exponențiale, despre care Kendall, Stuart [4], arată că sunt complete.

3 Estimatori liniari nedeplasați optimali (BLUE)

3.1 Introducere

În practică se întâmplă frecvent ca estimatorul MVU, dacă există, să nu poate fi găsit. De exemplu, este posibil să nu cunoaștem densitatea de repartiție a datelor sau chiar să putem fi dispuși să ne asumăm un model pentru acesta. În acest caz, metodele noastre anterioare, care se bazează pe CRLB și teoria statisticilor suficiente, nu pot fi aplicate. Chiar dacă densitatea este cunoscută, ultimele abordări nu garantează obținerea estimatorului

MVU. Confruntându-ne cu incapacitatea de a determina estimatorul *MVU* optimal, este rezonabil să apelăm la un estimator suboptimal. Procedând astfel, nu putem fi siguri de cât de multă performanță am putea pierde (deoarece dispersia minimală a estimatorului *MVU* este necunoscută). Cu toate acestea, dacă dispersia estimatorului suboptimal poate fi stabilită și dacă îndeplinește specificațiile sistemului nostru, atunci utilizarea sa poate fi justificată ca fiind adecvată pentru problema în cauză. Dacă dispersia sa este prea mare, atunci va trebui să ne uităm la alți estimatori suboptimali, sperând să găsim unul care să îndeplinească specificațiile noastre. O abordare frecventă este aceea de a restricționa estimatorul să fie liniar în raport cu datele și de a găsi estimatorul liniar care este nedeplasat și are dispersia minimă. După cum vom vedea, acest estimator, care este denumit *estimator liniar nedeplasat optimal* (*BLUE, the best linear unbiased estimator*), poate fi determinat doar cu cunoașterea momentelor de ordinul întâi și al doilea. Deoarece nu este necesară cunoașterea completă a densității de repartiție, *BLUE* este, de multe ori, cel mai potrivit pentru implementarea practică.

BLUE se bazează pe estimarea liniară definită de (56). Dacă impunem ca această estimare să fie nedeplasată, precum în (57), și să minimizăm dispersia lui (58), atunci *BLUE* este dat de formula (60), iar dispersia minimală a acestui estimator este dată de aceeași formulă. Pentru aflarea estimatorului, este nevoie doar de cunoașterea mediei și dispersiei datelor. Atât în cazul scalar, cât și în cel vectorial, dacă caracteristica este de tip Gaussian, atunci *BLUE* coincide cu *MVU*.

3.2 Modul de definire al BLUE

Considerăm o selecție repetată de volum n asupra unei caracteristici a cărei densitate de repartiție parametrizată este $f(x, \theta)$, vectorul empiric de selecție fiind $v = \{x_1, x_2, \dots, x_n\}$. *BLUE* solicită restricția ca estimatorul să fie de tip liniar. Mai precis, considerând $X_1, \dots, X_n$ variabilele de selecție corespunzătoare eșantionării,

$$\hat{\theta} = \sum_{i=0}^{n-1} a_i X_i, \quad (56)$$

unde a_i sunt coeficienți constanți, ce trebuie determinați. În funcție de alegerea acestora, putem genera diferiți estimatori pentru θ . Cu toate acestea, cel mai bun estimator (*BLUE*) este definit ca fiind acela care este nedeplasat și care are dispersia minimă. Deoarece ne limităm la clasa estimatorilor liniari, *BLUE* va fi optimal (cu alte cuvinte, este estimator *MVU*) doar atunci când *MVU* este, la rândul său, liniar. De exemplu, pentru problema estimării valorii mărimii curentului continuu în WGN (vezi Exemplul 1.3), estimatorul *MVU* este media de selecție:

$$\hat{\theta} = \bar{X} = \sum_{i=0}^{n-1} \frac{1}{n} X_i,$$

care este, evident, de tip liniar. Prin urmare, dacă ne restricționăm doar la clasa estimatorilor ce au această proprietate, nu vom pierde nimic deoarece estimatorul *MVU* aparține acestei clase. Pe de altă parte, pentru problema estimării valorii mediei a zgomotului repartizat uniform, estimatorul *MVU* este:

$$\hat{\theta} = \frac{n+1}{2n} \max_{i \in \{0, 1, \dots, n-1\}} X_i,$$

care este neliniar în raport cu datele. Dacă am restricționa căutarea doar la cazul estimatorilor liniari, atunci *BLUE* ar trebui să fie media de selecție, așa cum vom vedea în această secțiune. După cum este prezentat în Exemplul 2.8, diferența în performanță este substanțială. Din nefericire, fără cunoaștere densității de repartiție a statisticii, nu putem determina pierderea de performanță prin apelarea la *BLUE*. Pentru anumite probleme de estimare, utilizarea lui *BLUE* poate fi total nepotrivită. Dacă privim la problema estimării puterii unui WGN, se poate arăta ușor că estimatorul *MVU* este (vezi Exemplul 1.5):

$$\widehat{\sigma^2} = \frac{1}{n} \sum_{i=0}^{n-1} X_i^2,$$

care este neliniar. Dacă am forța ca estimatorul să fie liniar, similar formulei (56), $\widehat{\sigma^2} = \sum_{i=1}^n a_i X_i$, iar valoarea sa medie este:

$$\mathbb{E}(\widehat{\sigma^2}) = \sum_{i=0}^{n-1} a_i \mathbb{E}(X_i) = 0, \quad \text{deoarece} \quad \mathbb{E}(X_i) = 0, \quad \text{pentru fiecare } i \in \{1, 2, \dots, n\}.$$

Nu putem deci găsi niciun estimator liniar care să fie nedeplasat, cu atât mai puțin unul de dispersie minimă. Deși *BLUE* nu este potrivit acestei probleme, un *BLUE* ce utilizează datele modificate $Y_i = X_i^2$ va produce un

estimator viabil, deoarece pentru

$$\widehat{\sigma^2} = \sum_{i=0}^{n-1} a_i Y_i = \sum_{i=0}^{n-1} a_i X_i^2, \quad \text{avem} \quad \mathbb{E}(\widehat{\sigma^2}) = \sum_{i=0}^{n-1} a_i \sigma^2 = \sigma^2.$$

Există multe valori ale coeficienților a_i care ar satisface această restricție. Problema este care valori trebuie utilizate pentru a obține BLUE?

3.3 Determinarea lui BLUE

Pentru a determina BLUE, constrângem estimatorul $\hat{\theta}$ să fie liniar și nedeplasat, iar apoi determinăm coeficienții a_i care duc la minimizarea dispersiei. Avem, astfel,

$$\mathbb{E}(\hat{\theta}) = \sum_{i=0}^{n-1} a_i \mathbb{E}(X_i) = \theta. \quad (57)$$

Dispersia estimatorului $\hat{\theta}$ este

$$D^2(\hat{\theta}) = \mathbb{E} \left(\left(\sum_{i=0}^{n-1} a_i X_i - \mathbb{E} \left(\sum_{i=0}^{n-1} a_i X_i \right) \right)^2 \right).$$

Folosind (57) și utilizând abordarea vectorială $\mathbf{a} = (a_1, a_2, \dots, a_n)^t$, obținem:

$$D^2(\hat{\theta}) = \mathbb{E} \left((\mathbf{a}^t \mathbf{X} - \mathbf{a}^t \mathbb{E}(\mathbf{X}))^2 \right) = \mathbb{E} \left((\mathbf{a}^t (\mathbf{X} - \mathbb{E}(\mathbf{X})))^2 \right) = \mathbb{E} \left(\mathbf{a}^t (\mathbf{X} - \mathbb{E}(\mathbf{X})) (\mathbf{X} - \mathbb{E}(\mathbf{X}))^t \mathbf{a} \right) = \mathbf{a}^t C \mathbf{a}. \quad (58)$$

Vectorul ponderilor $\mathbf{a}$ se obține prin minimizarea (58), în raport cu restricția (57). Înainte de a continua, trebuie să presupunem o anumită formă a mediei $\mathbb{E}(X_i)$. Pentru a satisface condiția de nedeplasare, $\mathbb{E}(X_i)$ ar trebui să fie liniar în raport cu θ , adică

$$\mathbb{E}(X_i) = s_i \theta, \quad i \in \{0, 1, \dots, n-1\}, \quad (59)$$

unde coeficienții s_i sunt cunoscuți. Altfel, ar putea fi imposibil ca restricția să fie satisfăcută. De exemplu, dacă $\mathbb{E}(X_i) = \cos \theta$, atunci condiția de nedeplasare devine:

$$\sum_{i=0}^{n-1} a_i \cos \theta = \theta, \quad \text{pentru toți } \theta.$$

Pentru acest exemplu este clar că nu există coeficienții a_i care să satisfacă egalitatea anterioară. Observăm că, dacă scriem pe X_i sub forma $X_i = \mathbb{E}(X_i) + (X_i - \mathbb{E}(X_i))$, atunci, notând $W_i := X_i - \mathbb{E}(X_i)$ ca fiind zgomotul pe sistem,

$$X_i = \theta s_i + W_i, \quad \text{pentru toți } i \in \{1, 2, \dots, n\}.$$

Pentru a generaliza analiza, avem nevoie de o transformare neliniară a datelor observate, așa cum am descris deja. Cu ipoteza dată de (59), problema estimării se pune în modul următor. Pentru a determina BLUE, trebuie să minimizăm dispersia $D^2(\hat{\theta}) = \mathbf{a}^t C \mathbf{a}$, cu restricția de nedeplasare, care, datorită formulelor (57) și (59) devine:

$$\sum_{i=0}^{n-1} a_i \mathbb{E}(X_i) = \theta \iff \sum_{i=0}^{n-1} a_i s_i \theta = \theta \iff \sum_{i=0}^{n-1} a_i s_i = 1 \iff \mathbf{a}^t \mathbf{s} = 1,$$

unde $\mathbf{s} := (s_0, s_1, \dots, s_{n-1})^t$. Vom vedea că soluția acestei probleme de minimizare este:

$$\mathbf{a}_{opt} = \frac{C^{-1} \mathbf{s}}{\mathbf{s}^t C^{-1} \mathbf{s}},$$

iar BLUE are reprezentarea

$$\hat{\theta} = \frac{\mathbf{s}^t C^{-1} \mathbf{X}}{\mathbf{s}^t C^{-1} \mathbf{s}}, \quad \text{dispersia sa fiind dată de} \quad D^2(\hat{\theta}) = \frac{1}{\mathbf{s}^t C^{-1} \mathbf{s}}. \quad (60)$$

În conformitate cu (59), deoarece $\mathbb{E}(\mathbf{X}) = \theta \mathbf{s}$, BLUE este nedeplasat:

$$\mathbb{E}(\hat{\theta}) = \frac{\mathbf{s}^t C^{-1} \mathbb{E}(\mathbf{X})}{\mathbf{s}^t C^{-1} \mathbf{s}} = \frac{\mathbf{s}^t C^{-1} \theta \mathbf{s}}{\mathbf{s}^t C^{-1} \mathbf{s}} = \theta.$$

Pentru a determina BLUE, trebuie să cunoaștem doar pe $\mathbf{s}$ și C (matricea de covarianță) sau primele două momente, dar nu și întreaga densitate de repartiție, după cum vom vedea în exemplele prezentate în continuare.

Pentru moment, prezentăm modul de obținere a valorii $\mathbf{a}_{opt}$, rezultată în urma rezolvării unei probleme de optimizare:

$$\begin{cases} \min D^2(\hat{\theta}) = \min \mathbf{a}^t C \mathbf{a} \\ \mathbf{a}^t \mathbf{s} = 1. \end{cases}$$

Vom folosi metoda multiplicatorilor lui Lagrange. Funcția Lagrangiană este $J = \mathbf{a}^t C \mathbf{a} + \lambda (\mathbf{a}^t \mathbf{s} - 1)$, al cărei gradient, în raport cu vectorul $\mathbf{a}$, este $\partial J / \partial \mathbf{a} = 2C\mathbf{a} + \lambda \mathbf{s}$. Egalându-l cu zero, obținem:

$$\mathbf{a} = -\frac{\lambda}{2} C^{-1} \mathbf{s}, \quad \text{iar} \quad \mathbf{a}^t \mathbf{s} = -\frac{\lambda}{2} \mathbf{s}^t C^{-1} \mathbf{s} = 1 \implies -\frac{\lambda}{2} = \frac{1}{\mathbf{s}^t C^{-1} \mathbf{s}}.$$

Rezultă deci că

$$\mathbf{a}_{opt} = \frac{C^{-1} \mathbf{s}}{\mathbf{s}^t C^{-1} \mathbf{s}} \quad \text{și} \quad D^2(\hat{\theta}) = \mathbf{a}_{opt}^t C \mathbf{a}_{opt} = \frac{\mathbf{s}^t C^{-1} C C^{-1} \mathbf{s}}{(\mathbf{s}^t C^{-1} \mathbf{s})^2} = \frac{1}{\mathbf{s}^t C^{-1} \mathbf{s}}.$$

Pentru a vedea faptul că $\mathbf{a}_{opt}$ este minimul global și nu doar un minim local considerăm funcția:

$$\begin{aligned} G &= (\mathbf{a} - \mathbf{a}_{opt})^t C (\mathbf{a} - \mathbf{a}_{opt}) = \mathbf{a}^t C \mathbf{a} - 2\mathbf{a}_{opt}^t C \mathbf{a} + \mathbf{a}_{opt}^t C \mathbf{a}_{opt} \\ &= \mathbf{a}^t C \mathbf{a} - 2 \frac{\mathbf{s}^t C^{-1} C \mathbf{a}}{\mathbf{s}^t C^{-1} \mathbf{s}} + \frac{1}{\mathbf{s}^t C^{-1} \mathbf{s}} = \mathbf{a}^t C \mathbf{a} - \frac{1}{\mathbf{s}^t C^{-1} \mathbf{s}}. \end{aligned}$$

Prin urmare,

$$\mathbf{a}^t C \mathbf{a} = (\mathbf{a} - \mathbf{a}_{opt})^t C (\mathbf{a} - \mathbf{a}_{opt}) + \frac{1}{\mathbf{s}^t C^{-1} \mathbf{s}},$$

care este minimizată doar de $\mathbf{a} = \mathbf{a}_{opt}$. Aceasta se întâmplă deoarece proprietatea de pozitivă definire a matricei C face ca primul termen din membrul din dreapta al egalității anterioare să fie strict mai mare decât zero pentru $\mathbf{a} - \mathbf{a}_{opt} \neq \mathbf{0}$ (vectorul nul).

Exemplul 3.1 Mărimea curentului continuu în WGN.

(a) Presupunem că datele observate sunt date de $X_i = A + W_i$, $i \in \{0, 1, \dots, n-1\}$, unde, pentru fiecare i , W_i este un zgomot alb, având dispersia σ^2 (și densitatea de repartiție necunoscută). Problema constă în estimarea mărimii A a curentului. Cum W_i nu sunt, neapărat, de tip Gaussian, zgomotele pot fi statistic dependente, chiar dacă ele sunt, două câte două, necorelate. Deoarece $\mathbb{E}(X_i) = A$, atunci, conform relației (59), $s_i = 1$ și, prin urmare, $\mathbf{s} = \mathbf{1}$. Reprezentarea (60) furnizează forma lui BLUE:

$$\hat{A} = \frac{\mathbf{1}^t \frac{1}{\sigma^2} I \mathbf{X}}{\mathbf{1}^t \frac{1}{\sigma^2} I \mathbf{1}} = \frac{1}{n} \sum_{i=0}^{n-1} X_i = \bar{X}, \quad \text{având dispersia} \quad D^2(\hat{A}) = \frac{1}{\mathbf{1}^t \frac{1}{\sigma^2} \mathbf{1}} = \frac{\sigma^2}{n}.$$

Prin urmare, media de selecție este BLUE, afirmație ce este adevărată independent de densitatea de repartiție a caracteristicii. În plus, este estimatorul de dispersie minimă pentru zgomotul Gaussian.

(b) Presupunem acum că toate W_i , sunt zgomote albe necorelate, de medie zero și cu dispersia $D^2(W_i) = \sigma_i^2$, $i \in \{0, 1, \dots, n-1\}$. La fel ca și la punctul precedent, $\mathbf{s} = \mathbf{1}$ și, conform cu (60),

$$\hat{A} = \frac{\mathbf{1}^t C^{-1} \mathbf{X}}{\mathbf{1}^t C^{-1} \mathbf{1}}, \quad \text{cu dispersia} \quad D^2(\hat{\theta}) = \frac{1}{\mathbf{1}^t C^{-1} \mathbf{1}}.$$

Matricea de covarianță, respectiv inversa sa, sunt date de formulele:

$$C = \begin{pmatrix} \sigma_1^2 & 0 & \dots & 0 \\ 0 & \sigma_2^2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \sigma_n^2 \end{pmatrix} \quad \text{și} \quad C^{-1} = \begin{pmatrix} \frac{1}{\sigma_1^2} & 0 & \dots & 0 \\ 0 & \frac{1}{\sigma_2^2} & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \frac{1}{\sigma_n^2} \end{pmatrix}.$$

Prin urmare,

$$\hat{A} = \frac{\sum_{i=0}^{n-1} \frac{X_i}{\sigma_i^2}}{\sum_{i=0}^{n-1} \frac{1}{\sigma_i^2}} \quad \text{și} \quad D^2(\hat{A}) = \frac{1}{\sum_{i=0}^{n-1} \frac{1}{\sigma_i^2}}. \quad (61)$$

În concluzie, BLUE ponderează mai mult acele probe ce au dispersia mai mică, în încercarea de a egaliza contribuția zgomotelor provenite de la fiecare probă în parte. Numitorul din formulele (61) este factorul de ajustare ce face ca estimatorul obținut să fie nedepășat.

Datorită ipotezei zgomotului Gaussian, estimatorul poate fi considerat a fi eficient și, prin urmare este un estimator *MVU*. Nu este doar o coincidență faptul că estimatorul *MVU* pentru cazul zgomotului Gaussian și *BLUE* sunt aceeași. Avem de a face, de fapt, cu un rezultat general, care afirmă că, în estimarea parametrilor unui model liniar, cele două tipuri de estimatori coincid dacă ne situăm în ipoteza menționată privind tipul zgomotului pe sistem.

3.4 Cazul parametrului multidimensional

Dacă parametrul estimat este un vector $p \times 1$ dimensional, atunci, pentru ca estimatorul să fie liniar în raport cu datele ar trebui ca el să aibă forma:

$$\hat{\theta}_i = \sum_{i=0}^{n-1} a_{in} X_i, \quad i \in \{1, 2, \dots, p\}, \quad (62)$$

unde a_{in} sunt coeficienții de ponderare. Scrisă sub formă matriceală, formula (62) devine $\hat{\theta} = \mathbf{A}V$, unde $\mathbf{A} \in \mathcal{M}_{p \times n}(\mathbb{R})$. Pentru ca estimatorul $\hat{\theta}$ să fie nedepășat, trebuie ca

$$\mathbb{E}(\hat{\theta}_i) = \sum_{i=0}^{n-1} a_{in} \mathbb{E}(X_i) = \theta_i, \quad i \in \{1, 2, \dots, p\}, \quad \text{sau, echivalent,} \quad \mathbb{E}(\hat{\theta}) = \mathbf{A}\mathbb{E}(V) = \theta. \quad (63)$$

Un estimator liniar este potrivit doar pentru probleme în care restricția privind nedepășarea poate fi satisfăcută. Din formula matriceală precedentă obținem că:

$$\mathbb{E}(V) = \mathbf{H}\theta, \quad \text{unde } \mathbf{H} \in \mathcal{M}_{n \times p}(\mathbb{R}). \quad (64)$$

În cazul parametrului scalar (vezi (59)), $\mathbb{E}(V) = \mathbf{H}\theta = (s_0, s_1, \dots, s_{n-1})^t \theta$. Substituind (64) în (63), obținem condiția de nedepășare

$$\mathbf{A}\mathbf{H} = \mathbf{I}. \quad (65)$$

Dacă definim linia $\mathbf{a}_i := (a_{i0}, a_{i1}, \dots, a_{i(n-1)})^t$, astfel încât $\hat{\theta}_i = \mathbf{a}_i^t V$, atunci condiția de nedepășare poate fi rescrisă pentru fiecare $\mathbf{a}_i$ notând

$$\mathbf{A} = \begin{pmatrix} \mathbf{a}_1^t \\ \mathbf{a}_2^t \\ \vdots \\ \mathbf{a}_p^t \end{pmatrix} \quad \text{și considerând } \mathbf{h}_i \text{ coloana } i \text{ a lui } \mathbf{H} = \begin{pmatrix} \mathbf{h}_1 & \mathbf{h}_2 & \cdots & \mathbf{h}_p \end{pmatrix}.$$

Astfel, condiția de nedepășare (65) se reduce la

$$\mathbf{a}_i^t \mathbf{h}_j = \delta_{ij}, \quad i, j \in \{1, 2, \dots, p\}. \quad (66)$$

Dispersia estimatorilor va fi, pentru fiecare i ,

$$D^2(\hat{\theta}_i) = \mathbf{a}_i^t C \mathbf{a}_i. \quad (67)$$

BLUE pentru parametrul vectorial se obține prin minimizarea cantităților (67), cu restricția dată de (66), pentru fiecare i . Arătăm în continuare că *BLUE* și matricea de covarianță sunt date de formulele:

$$\hat{\theta} = (\mathbf{H}^t C^{-1} \mathbf{H})^{-1} \mathbf{H}^t C^{-1} V, \quad \text{iar} \quad C_{\hat{\theta}} = (\mathbf{H}^t C^{-1} \mathbf{H})^{-1}. \quad (68)$$

Ne situăm în ipoteza unui model liniar, $V = \mathbf{H}\theta + \mathbf{w}$, cu $\mathbf{w} \sim \mathcal{N}(\mathbf{0}, C)$. Scopul este deci acela de a estima parametrul vectorial θ , unde $\mathbb{E}(V) = \mathbf{H}\theta$. Aceasta este exact presupunerea dată de formula (64). Dar, estimatorul *MVU* pentru date de tip Gaussian este

$$\hat{\theta} = (\mathbf{H}^t C^{-1} \mathbf{H})^{-1} \mathbf{H}^t C^{-1} V,$$

care este liniar în V . Putem concluziona că dacă datele sunt de tip Gaussian, atunci *BLUE* este și estimator *MVU*.

Problema minimizării dispersiilor date de (67), în prezența restricțiilor de nedepășire (66) este similară cazului scalar, exceptând prezența restricțiilor asupra vectorilor $\mathbf{a}_i$. Avem acum p restricții asupra acestor vectori, nu doar una. Cum vectorii $\mathbf{a}_i$ pot lua orice valori, independent între ei, avem de a face, de fapt, cu p probleme de minimizare separate, interconectate doar de restricții. Considerăm, pentru fiecare i , Lagrangian-ul

$$J_i = \mathbf{a}_i^t C \mathbf{a}_i + \sum_{j=1}^p \lambda_j^{(i)} (\mathbf{a}_i^t \mathbf{h}_j - \delta_{ij}), \quad \text{cu derivatele parțiale} \quad \frac{\partial J_i}{\partial \mathbf{a}_i} = 2C \mathbf{a}_i + \sum_{j=1}^p \lambda_j^{(i)} \mathbf{h}_j = 2C \mathbf{a}_i + \mathbf{H} \lambda_i,$$

pentru ultima egalitate folosind notațiile $\lambda_i := (\lambda_1^{(i)}, \lambda_2^{(i)}, \dots, \lambda_p^{(i)})^t$ și $\mathbf{H} = (\mathbf{h}_1 \quad \mathbf{h}_2 \quad \dots \quad \mathbf{h}_p)$. Egalând gradientul cu zero, găsim:

$$\mathbf{a}_i = -\frac{1}{2} C^{-1} \mathbf{H} \lambda_i. \quad (69)$$

Pentru a determina multiplicatorii lui Lagrange, ne vom folosi de restricțiile (66) care, în formă combinată, duc la ecuația:

$$\begin{pmatrix} \mathbf{h}_1^t \\ \vdots \\ \mathbf{h}_{i-1}^t \\ \mathbf{h}_i^t \\ \mathbf{h}_{i+1}^t \\ \vdots \\ \mathbf{h}_p^t \end{pmatrix} \mathbf{a}_i = \begin{pmatrix} 0 \\ \vdots \\ 0 \\ 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix}.$$

Notăm cu e_i vectorul canonic i și obținem restricția $\mathbf{H}^t \mathbf{a}_i = e_i$. Folosim (69), iar restricția devine:

$$\mathbf{H}^t \mathbf{a}_i = -\frac{1}{2} \mathbf{H}^t C^{-1} \mathbf{H} \lambda_i = e_i.$$

Dacă presupunem inversabilitatea matricei $\mathbf{H}^t C^{-1} \mathbf{H}$, atunci vectorul multiplicatorilor lui Lagrange este:

$$-\frac{1}{2} \lambda_i = (\mathbf{H}^t C^{-1} \mathbf{H})^{-1} e_i, \quad \text{iar} \quad \mathbf{a}_{i_{opt}} = C^{-1} \mathbf{H} (\mathbf{H}^t C^{-1} \mathbf{H})^{-1} e_i. \quad (70)$$

Dispersia acestui estimator este:

$$D^2(\hat{\theta}_i) = \mathbf{a}_{i_{opt}}^t C \mathbf{a}_{i_{opt}} = e_i^t (\mathbf{H}^t C^{-1} \mathbf{H})^{-1} \mathbf{H}^t C^{-1} C C^{-1} \mathbf{H} (\mathbf{H}^t C^{-1} \mathbf{H})^{-1} e_i = e_i^t (\mathbf{H}^t C^{-1} \mathbf{H})^{-1} e_i.$$

Similar cazului scalar, se poate arăta că $\mathbf{a}_{i_{opt}}$ este un punct de minim global. Putem exprima vectorul BLUE sub o formă mai compactă, astfel:

$$\hat{\theta} = \begin{pmatrix} \mathbf{a}_{1_{opt}}^t V \\ \mathbf{a}_{2_{opt}}^t V \\ \vdots \\ \mathbf{a}_{p_{opt}}^t V \end{pmatrix} = \begin{pmatrix} e_1^t (\mathbf{H}^t C^{-1} \mathbf{H})^{-1} \mathbf{H}^t C^{-1} V \\ e_2^t (\mathbf{H}^t C^{-1} \mathbf{H})^{-1} \mathbf{H}^t C^{-1} V \\ \vdots \\ e_p^t (\mathbf{H}^t C^{-1} \mathbf{H})^{-1} \mathbf{H}^t C^{-1} V \end{pmatrix} = \begin{pmatrix} e_1^t \\ e_2^t \\ \vdots \\ e_p^t \end{pmatrix} (\mathbf{H}^t C^{-1} \mathbf{H})^{-1} \mathbf{H}^t C^{-1} V = (\mathbf{H}^t C^{-1} \mathbf{H})^{-1} \mathbf{H}^t C^{-1} V,$$

deoarece $(e_1^t \ e_2^t \dots \ e_p^t)$ este matricea unitate. De asemenea, matricea de covarianță a estimatorului vectorial $\hat{\theta}$ este:

$$C_{\hat{\theta}} = \mathbb{E} \left((\hat{\theta} - \mathbb{E}(\hat{\theta})) (\hat{\theta} - \mathbb{E}(\hat{\theta}))^t \right),$$

unde

$$\hat{\theta} - \mathbb{E}(\hat{\theta}) = (\mathbf{H}^t C^{-1} \mathbf{H})^{-1} \mathbf{H}^t C^{-1} (\mathbf{H} \theta + \mathbf{w}) - \mathbb{E}(\hat{\theta}) = (\mathbf{H}^t C^{-1} \mathbf{H})^{-1} \mathbf{H}^t C^{-1} \mathbf{w}.$$

Obținem astfel:

$$\begin{aligned} C_{\hat{\theta}} &= \mathbb{E} \left((\mathbf{H}^t C^{-1} \mathbf{H})^{-1} \mathbf{H}^t C^{-1} \mathbf{w} \mathbf{w}^t C^{-1} \mathbf{H} (\mathbf{H}^t C^{-1} \mathbf{H})^{-1} \right) \\ &= (\mathbf{H}^t C^{-1} \mathbf{H})^{-1} \mathbf{H}^t C^{-1} C C^{-1} \mathbf{H} (\mathbf{H}^t C^{-1} \mathbf{H})^{-1} = (\mathbf{H}^t C^{-1} \mathbf{H})^{-1}, \end{aligned}$$

deoarece matricea de covarianță a lui $\mathbf{w}$ este C . Minimul dispersiilor este dat de către elementele diagonale ale matricei $C_{\hat{\theta}}$, adică:

$$D^2(\hat{\theta}_i) = e_i^t C_{\hat{\theta}} e_i = (\mathbf{H}^t C^{-1} \mathbf{H})_{ii}^{-1}, \quad i \in \{1, 2, \dots, p\},$$

obținând astfel rezultatul dorit.

În cadrul considerat, modelul liniar general nu presupune existența zgomotului Gaussian. Datele au doar o formă liniară de reprezentare. Următorul rezultat oferă formula estimatorului și a matricei sale de covarianță în această situație.

Teorema 6 (Gauss-Markov) Dacă datele au forma liniară generală:

$$V = \mathbf{H}\theta + \mathbf{w}, \quad \text{unde } \mathbf{H} \in \mathcal{M}_{n \times p}(\mathbb{R}) \text{ este cunoscută, } \theta \in \mathcal{M}_{p \times 1} \text{ este vectorul parametrilor,}$$

iar $\mathbf{w}$ este vectorul zgomotelor pe sistem, de tip $n \times 1$, având media zero și matricea de covarianță C (densitatea de repartiție fiind necunoscută, nu neapărat de tip Gaussian), atunci BLUE pentru θ este:

$$\hat{\theta} = (\mathbf{H}^t C^{-1} \mathbf{H})^{-1} \mathbf{H}^t C^{-1} V, \quad \text{iar dispersia minimă pentru } \hat{\theta}_i \text{ este } D^2(\hat{\theta}_i) = (\mathbf{H}^t C^{-1} \mathbf{H})_{ii}^{-1}, \quad i \in \{1, 2, \dots, p\}. \quad (71)$$

În plus, matricea de covarianță a estimatorului vectorial $\hat{\theta}$ este:

$$C_{\hat{\theta}} = (\mathbf{H}^t C^{-1} \mathbf{H})^{-1}.$$

3.5 Un exemplu privind procesarea semnalului. Localizarea sursei

O problemă de interes privind traficul aerian îl constituie determinarea poziției unei surse emitente a unui semnal. De exemplu, pentru a determina poziția unui avion, se utilizează, de regulă, o rețea de antene. Metoda constă în estimarea diferențelor temporale dintre o serie de măsurători (TDOA, *time difference of arrival*). Astfel, semnalul emis de un avion (sursa) este recepționat, în același timp, de trei antene din rețea, având deci $t_1 = t_2 = t_3$. TDOA $t_2 - t_1$ dintre antena 1 și 2 este zero, la fel ca și $t_3 - t_2$ dintre antenele 2 și 3. Putem afirma astfel că range-urile (razele de acțiune, distanțele) R_i către cele trei antene sunt aceleași. Intersecția porțiunilor de disc având ca centru antenele și ca rază range-urile R_i , reprezintă posibila locație a avionului emitent al semnalului.

Vom examina problema de localizare folosind teoria estimației. Considerăm o rețea de n antene ce sunt poziționate în locații cunoscute și că măsurătorile timpilor de sosire (recepționare a semnalului de către antene) t_i , $i \in \{0, 1, \dots, n-1\}$ sunt disponibile. Trebuie să estimăm poziția sursei (x_s, y_s) . De asemenea, presupunem că timpii de sosire sunt perturbați de câte un zgomot aleator de medie zero, dispersie cunoscută, dar densitate de repartiție necunoscută. Pentru semnalul emis de sursă la momentul $t = T_0$, măsurătorile sunt modelate de relațiile:

$$t_i = T_0 + R_i/c + \varepsilon_i, \quad i \in \{0, 1, \dots, n-1\}, \quad (72)$$

unde ε_i sunt zgomotele (perturbațiile) măsurătorilor, iar c reprezintă viteza de propagare. Zgomotele măsurătorilor sunt de medie 0, dispersie σ^2 și necorelate între ele. Presupunem că nu există nicio informație privind distribuția zgomotelor. Vom stabili o relație între raza de acțiune R_i a fiecărei antene (distanța până la țintă) și poziția necunoscută $\theta = \begin{pmatrix} x_s & y_s \end{pmatrix}^t$ a emițătorului. Notăm poziția fiecărei antene cu $(x_i, y_i)_i$, poziții pe care le presupunem cunoscute și avem:

$$R_i = \sqrt{(x_s - x_i)^2 + (y_s - y_i)^2} \quad (73)$$

Dacă introducem formula lui R_i din (73) în (72), se observă că modelul este neliniar în raport cu parametrii necunoscuți x_s și y_s . Pentru a aplica Teorema Gauss-Markov, presupunem că este disponibilă o poziție probabilă a sursei, (x_m, y_m) . Această poziție probabilă, care este apropiată de adevărata poziție a sursei (poziția exactă) poate fi determinată din măsurători anterioare. O astfel de situație apare, de regulă, când sursa este urmărită de receptori. În consecință, pentru a estima poziția nouă a sursei, trebuie să estimăm parametrul $\theta = \begin{pmatrix} x_s - x_m & y_s - y_m \end{pmatrix}^t = \begin{pmatrix} \delta x_s & \delta y_s \end{pmatrix}^t$. Fie R_{m_i} razele de acțiune nominale, precum în Figura 5.

Folosim o dezvoltare Taylor de ordinul întâi a cantității R_i (considerată ca o funcție de două variabile, x_s și y_s), în jurul poziției nominale $x_s = x_m, y_s = y_m$:

$$R_i \approx R_{m_i} + \frac{x_m - x_i}{R_{m_i}} \delta x_s + \frac{y_m - y_i}{R_{m_i}} \delta y_s. \quad (74)$$

Folosim această aproximare în formula (72) și obținem:

$$t_i = T_0 + \frac{R_{m_i}}{c} + \frac{x_m - x_i}{R_{m_i} c} \delta x_s + \frac{y_m - y_i}{R_{m_i} c} \delta y_s + \varepsilon_i, \quad i \in \{0, 1, \dots, n-1\},$$

formulă care este, acum, liniară în raport cu parametrii δx_s și δy_s .

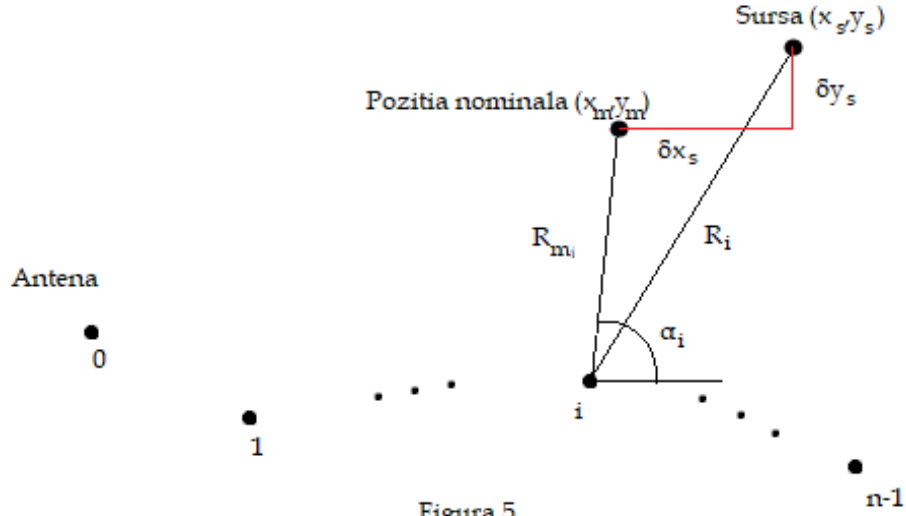


Figura 5.

De asemenea, deoarece

$$\cos \alpha_i = \frac{x_m - x_i}{R_{m_i}} \quad \text{și} \quad \sin \alpha_i = \frac{y_m - y_i}{R_{m_i}},$$

modelul se poate scrie, simplificat,

$$t_i = T_0 + \frac{R_{m_i}}{c} + \frac{\cos \alpha_i}{c} \delta x_s + \frac{\sin \alpha_i}{c} \delta y_s + \varepsilon_i, \quad i \in \{0, 1, \dots, n-1\}.$$

Termenul R_{m_i}/c este o constantă cunoscută care poate fi încorporată în măsurătoare prin notația

$$\tau_i = t_i - \frac{R_{m_i}}{c}.$$

Modelul nostru linear devine:

$$\tau_i = T_0 + \frac{\cos \alpha_i}{c} \delta x_s + \frac{\sin \alpha_i}{c} \delta y_s + \varepsilon_i, \quad i \in \{0, 1, \dots, n-1\}, \quad (75)$$

unde parametri necunoscuți sunt $T_0, \delta x_s, \delta y_s$. Presupunerea că momentul T_0 la care sursa emite un semnal este necunoscut este des întâlnită în cazurile concrete. Cunoașterea lui T_0 ar presupune o sincronizare precisă a ceasurilor sursei și receiver-ului, fapt care, în situații reale este de evitat. De regulă, se consideră *TDOA* care elimină pe T_0 din (75). Construim această diferență de timp *TDOA* în felul următor:

$$\begin{cases} \xi_1 &= \tau_1 - \tau_0 \\ \xi_2 &= \tau_2 - \tau_1 \\ &= \\ \xi_{n-1} &= \tau_{n-1} - \tau_{n-2}. \end{cases}$$

Formula (75) conduce la următorul model (final) liniar:

$$\xi_i = \frac{1}{c} (\cos \alpha_i - \cos \alpha_{i-1}) \delta x_s + \frac{1}{c} (\sin \alpha_i - \sin \alpha_{i-1}) \delta y_s + \varepsilon_i - \varepsilon_{i-1}, \quad i \in \{1, 2, \dots, n-1\}. \quad (76)$$

Am redus problema estimării la cea descrisă de Teorema Gauss-Markov, unde

$$\theta = \begin{pmatrix} \delta x_s \\ \delta y_s \end{pmatrix}, \quad \mathbf{H} = \frac{1}{c} \begin{pmatrix} \cos \alpha_1 - \cos \alpha_0 & \sin \alpha_1 - \sin \alpha_0 \\ \cos \alpha_2 - \cos \alpha_1 & \sin \alpha_2 - \sin \alpha_1 \\ \vdots & \vdots \\ \cos \alpha_{n-1} - \cos \alpha_{n-2} & \sin \alpha_{n-1} - \sin \alpha_{n-2} \end{pmatrix}, \quad \mathbf{w} = \begin{pmatrix} \varepsilon_1 - \varepsilon_0 \\ \varepsilon_2 - \varepsilon_1 \\ \vdots \\ \varepsilon_{n-1} - \varepsilon_{n-2} \end{pmatrix}.$$

Vectorul zgomotelor $\mathbf{w}$ are media zero, dar nu mai este compus din variabile aleatoare necorelate între ele. Pentru a determina matricea de covarianță, observăm că:

$$\mathbf{w} = \mathbf{A}\varepsilon = \begin{pmatrix} -1 & 1 & 0 & 0 & \cdots & 0 \\ 0 & -1 & 1 & 0 & \cdots & 0 \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & \cdots & 0 & -1 & 1 \end{pmatrix} \begin{pmatrix} \varepsilon_0 \\ \varepsilon_1 \\ \vdots \\ \varepsilon_{n-1} \end{pmatrix}, \quad \mathbf{A} \in \mathcal{M}_{(n-1) \times n}(\mathbb{R}).$$

Deoarece matricea de covarianță a lui ε este $\sigma^2 I$, obținem:

$$C = \mathbb{E}(\mathbf{A}\varepsilon\varepsilon^t\mathbf{A}^t) = \sigma^2\mathbf{A}\mathbf{A}^t.$$

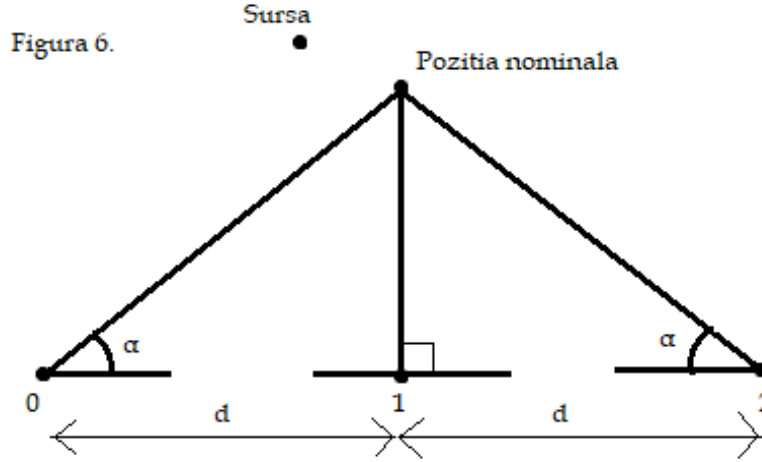
Din (71), BLUE pentru parametrul vectorial al poziției sursei este:

$$\hat{\theta} = (\mathbf{H}^t C^{-1} \mathbf{H})^{-1} \mathbf{H}^t C^{-1} \xi = (\mathbf{H}^t (\mathbf{A}\mathbf{A}^t) \mathbf{H})^{-1} \mathbf{H}^t (\mathbf{A}\mathbf{A}^t)^{-1} \xi, \quad (77)$$

iar dispersia minimă, respectiv matricea de covarianță, sunt date de:

$$D^2(\hat{\theta}_i) = \sigma^2 \left((\mathbf{H}^t (\mathbf{A}\mathbf{A}^t)^{-1} \mathbf{H})^{-1} \right)_{ii} \quad \text{și} \quad C_{\hat{\theta}} = \sigma^2 (\mathbf{H}^t (\mathbf{A}\mathbf{A}^t)^{-1} \mathbf{H})^{-1}. \quad (78)$$

Ca un exemplu, în Figura 6, avem o rețea de 3 antene în linie, care este, de fapt, numărul minim de antene necesare pentru un model de localizare.



Pentru acest model avem:

$$\mathbf{H} = \frac{1}{c} \begin{pmatrix} -\cos \alpha & 1 - \sin \alpha \\ -\cos \alpha & -(1 - \sin \alpha) \end{pmatrix}, \quad \mathbf{A} = \begin{pmatrix} -1 & 1 & 0 \\ 0 & -1 & 1 \end{pmatrix},$$

iar matricea de covarianță este, conform formulei (78),

$$C_{\hat{\theta}} = \sigma^2 c^2 \begin{pmatrix} \frac{1}{2 \cos^2 \alpha} & 0 \\ 0 & \frac{3/2}{(1 - \sin \alpha)^2} \end{pmatrix}.$$

Pentru o mai bună localizare, unghiul α ar trebui să fie cât mai mic. Aceasta se poate realiza mărinnd spațiul d dintre antene. Remarcăm, de asemenea, că acuratețea locației este dependentă de range, acuratețea localizării îmbunătățindu-se pentru range-uri scurte sau pentru unghiuri α mici.

Bibliografie

- [1] Devore, J; Berk, K., *Modern Mathematical Statistics with Applications*, 2nd Edition, Springer New York Dordrecht Heidelberg London, 2012.
- [2] Kay, S., *Fundamentals of Statistical Signal Processing, Volume I: Estimation Theory*, Prentice Hall; 1 edition, April 5, 1993.
- [3] Kay, S., *Fundamentals of Statistical Signal Processing, Volume II: Detection Theory*, Pearson; 1 edition, February 6, 1998.
- [4] Kendall, M.G.; Stuart, A., *The Advanced Theory of Statistics*, Volume 2, Inference and Relationships, Hafner Publishing Company, 1961 (Edition by Wiley, 2010).
- [5] Montgomery, D; Runger, G, *Applied Statistics and Probability for Engineers*, 3rd Edition, John Wiley & Sons, Inc, 2003.
- [6] Wackerly, D.; Mendenhall, W.; Scheaffer, R., *Mathematical Statistics with Applications*, 7th Edition, Thomson Brooks/Cole, 2008.